

INCLUVERSO 5G

Tecnologías de realidad exteNdida y Comunicaciones para la incLUsión de personas en situación de VulnERabilidad psicoSOcial mediante redes 5G avanzadas



Financiado por
la Unión Europea
NextGenerationEU



GOBIERNO
DE ESPAÑA

MINISTERIO
PARA LA TRANSFORMACIÓN DIGITAL
Y DE LA FUNCIÓN PÚBLICA



Plan de Recuperación,
Transformación
y Resiliencia



UNICO
I+D

Entregable E2.3

Tecnologías de redes 5G y telecomunicaciones

Contenidos

1	Introducción	5
2	Validación de la red de comunicaciones en banda n40 usada en el proyecto	6
	Descripción del despliegue en banda n40 en la Fundación Juan XXIII	7
	Pruebas de cobertura	8
	Pruebas de slicing	11
2.1	2.3.1 Configuración dinámica de slicing en uplink	11
2.2	2.3.2 Configuración dinámica de slicing en downlink	13
2.3	2.3.3 Pruebas en dos módems	14
	Pruebas de estabilidad	15
2.4	Conclusiones	21
2.5	Validación de los modelos y técnicas para modelar y controlar QoS/QoE.....	23
3.1	Modelos paramétricos de QoE para los escenarios de aplicación del proyecto	23
	3.1.1 Escenario 1: Teleportación Virtual (Virtual Teleportation)	24
	3.1.2 Escenario 2: Procesamiento delegado de Realidad Mixta (<i>Mixed Reality computation offloading</i>)	25
3.2	3.1.3 Escenario 3: Regulación Emocional (<i>Emotional Regulation</i>)	26
3.3	Validación de las mejoras en el emulador FikoRE	27
	Pruebas en red emulada (5QIs)	31
	3.3.1 Escenario 1: Teleportación Virtual (Virtual Teleportation)	33
3.4	3.3.2 Escenario 2: Procesamiento delegado de Realidad Mixta (<i>Mixed Reality computation offloading</i>)	35
	3.3.3 Escenario 3: Regulación Emocional (<i>Emotional Regulation</i>)	39
	Pruebas en red física	41
	3.4.1 Escenario 1: Teleportación Virtual (Virtual Teleportation)	42
	3.4.2 Escenario 2: Procesamiento delegado de Realidad Mixta (<i>Mixed Reality computation offloading</i>)	50
	3.4.3 Escenario 3: Regulación Emocional (<i>Emotional Regulation</i>)	53

	Conclusiones	56
4	Pruebas finales del sistema de edge y de offloading de IA	58
	Pruebas funcionales.....	58
3.5	4.1.1 Arquitectura general.....	58
	4.1.2 Sistema de virtualización	60
4.1	4.1.3 VNF de offloading de algoritmos de IA para XR.....	64
	4.1.4 VNF de backend de telepresencia.....	66
	4.1.5 VNF de backend de teleformación.....	67
	4.1.6 VNF de monitorización y gestión de KPIs	67
	Pruebas de carga.....	68
4.2	4.2.2 Escenario 3: Regulación Emocional (<i>Emotional Regulation</i>)	79
	4.2.3 Conclusiones	80
5	Evaluación del rendimiento de la plataforma desplegada en la Fundación Juan XXIII en los distintos los casos de uso	81
5.1	Terapia conductual	81
	5.1.1 Sistemas desplegados.....	81
	5.1.2 Evaluación del rendimiento.....	82
5.2	5.1.3 Discusión	82
	Telepresencia.....	83
	5.2.1 Sistemas desplegados.....	83
5.3	5.2.2 Evaluación del rendimiento.....	83
	5.2.3 Discusión	84
	Teleformación.....	85
	5.3.1 Sistemas desplegados.....	85
	5.3.2 Escenarios funcionales de teleformación.....	86
6.1	5.3.3 Evaluación del rendimiento.....	86
	5.3.4 Discusión	87
6	Apoyo de Fundación Juan XXIII	88
	Servicios en las instalaciones de la Fundación.....	88

6.1.1	Terapias conductuales o cognitivas	88
6.1.2	Estimulación multisensorial virtual	89
6.1.3	Entrenamiento en movilidad	90
6.1.4	Formación.....	91
6.1.5	Asistencia a eventos y formaciones	91
6.1.6	Vigilancia de exteriores	92
	Servicios en remoto	92
7	Conclusiones	94
8 6.2	Acrónimos.....	95
9	Índice de tablas.....	97
10	Índice de figuras	99
11	Referencias.....	104

1 Introducción

El presente entregable presenta la validación técnica de la infraestructura 5G y de la plataforma de realidad extendida (XR) desplegadas en el proyecto INCLUVERSO 5G. El objetivo principal del documento es evaluar la capacidad de la red 5G privada en banda n40 para dar soporte a servicios XR en un entorno sociosanitario real y operativo.

El análisis aborda de forma integrada los distintos componentes de la solución, incluyendo la conectividad 5G, las capacidades de computación en el borde (edge computing) y los mecanismos de control de calidad de servicio y calidad de experiencia (QoS/QoE). La validación se apoya en una combinación de pruebas realizadas tanto en red física como en red emulada, lo que permite caracterizar el comportamiento del sistema bajo condiciones representativas y controladas.

Se consideran tres escenarios base de aplicación: teleportación virtual, procesamiento delegado de realidad mixta y entrenamiento cognitivo/regulación emocional. El estudio se fundamenta en métricas objetivas de rendimiento de red y sistema, complementadas con modelos de calidad de experiencia que permiten relacionar dichos parámetros técnicos con la percepción del usuario final.

Finalmente, los resultados obtenidos se contextualizan en casos de uso reales desarrollados en la Fundación Juan XXIII, proporcionando una visión aplicada del potencial de la infraestructura desplegada y de su idoneidad para soportar servicios XR avanzados en entornos sociosanitarios.

2 Validación de la red de comunicaciones en banda n40 usada en el proyecto

Antes de abordar la validación de los modelos de calidad de servicio y calidad de experiencia, así como la evaluación de los distintos casos de uso del proyecto INCLUVERSO 5G, es necesario verificar el correcto funcionamiento de la infraestructura de comunicaciones desplegada en la Fundación Juan XXIII Roncalli. Este capítulo presenta la validación experimental de la red 5G privada en banda n40 utilizada como soporte de los servicios y aplicaciones desarrollados en el proyecto.

La red desplegada constituye una red 5G Standalone (SA) completa, en formato de “isla 5G”, que integra estación base, núcleo de red y capacidades de computación en el borde (MEC). Esta infraestructura se ha instalado en un entorno real y operativo, con múltiples espacios interiores y exteriores, y con una actividad diaria normal del centro, lo que permite evaluar su comportamiento en condiciones representativas de uso real.

El objetivo principal de este capítulo es comprobar que la red proporciona los niveles de cobertura, capacidad, estabilidad y control de calidad necesarios para dar soporte a los requisitos definidos en el entregable E1.2 para los distintos escenarios base del proyecto: teleportación virtual, procesamiento delegado de realidad mixta y regulación emocional. En particular, se analiza la capacidad de la red para ofrecer conectividad fiable en los espacios relevantes de la Fundación, así como la efectividad de los mecanismos de priorización de tráfico mediante network slicing, fundamentales para garantizar la calidad de los servicios críticos frente a tráfico competitivo.

Para ello, se han llevado a cabo distintas campañas de pruebas que incluyen medidas de cobertura radio, ensayos de control dinámico de capacidad en enlace ascendente y descendente, pruebas de coexistencia de múltiples usuarios y evaluación de la estabilidad del sistema bajo carga sostenida. Las métricas consideradas abarcan parámetros de rendimiento de red (throughput, latencia, pérdida de paquetes), comportamiento del núcleo y del nodo MEC, así como indicadores de calidad de señal radio.

Los resultados presentados en este capítulo permiten establecer una base sólida para el análisis posterior de la calidad de servicio y calidad de experiencia, y confirman la idoneidad de la infraestructura 5G desplegada para soportar aplicaciones XR y servicios avanzados en un entorno sociosanitario real.

Descripción del despliegue en banda n40 en la Fundación Juan XXIII

Se ha desplegado un nodo de red 5G en banda n40, en formato “isla 5G”, es decir, una red 5G completa autocontenida, incluyendo estación base (gNB), núcleo de red (“core”) y nodo de computación en el borde (MEC).

2.1 El gNB se compone de los siguientes elementos físicos:

- Banda base, instalada en bastidor, compuesta de los módulos hardware:
 - Nokia AirScale Subrack: chasis en el que se instala el módulo de banda base.
 - Nokia AirScale Common: elemento común de administración y alimentación de los módulos de banda base.
 - Nokia AirScale Capacity: módulo de banda base de 5G Standalone (SA).
- Módulo de antena, instalado en la cubierta del edificio sobre un mástil austosoportado. Incluye:
 - Cabeza radio remota (RRH) Nokia AirScale.
 - Antena omnidireccional.
 - Alimentador.
 - Receptor GPS

En el bastidor, además de la banda base del gNB, se ha instalado:

- MEC Nokia OpenEdge.
- Router Nokia IXR 7250
- Alimentador DC.

Se trata una única celda de red con una antena omnidireccional no sectorizada, con las siguientes características:

- Banda: n40
- Tecnología de celda: TDD
- Máxima potencia de salida: 46 dBm
- Ancho de banda: 20 MHz
- Frecuencia de portadora: 2380 MHz
- NRARFCN: 476000
- NR cell ID: 16394
- PLMN: 99 999

En la plataforma 5G que se ha desplegado se ha utilizado el core open5gs.

Se han configurado tres slices:

- Slice 1, PRIORITY (dnn Golden).
- Slice 2, HMDS (dnn Silver).
- Slice 3, DEFAULT (dnn Bronze).

El detalle del diseño se describe en el entregable E2.1. El despliegue de red se detalla en el E2.2.

Pruebas de cobertura

El objetivo de las pruebas de cobertura es evaluar la disponibilidad de conectividad 5G en las 2.2 distintas zonas de la Fundación Juan XXIII Roncalli en las que se desarrollan, o podrían desarrollarse, los casos de uso del proyecto. Estas pruebas permiten verificar que la red en banda n40 proporciona niveles de señal y capacidad suficientes tanto en interiores como en exteriores, y sirven como referencia para la planificación y validación de los escenarios posteriores.

Las medidas se realizaron en ubicaciones representativas del centro, incluyendo espacios de atención psicológica, aulas, talleres, zonas comunes, exteriores y áreas de transición. La selección de los puntos de medida, mostrada en la figura, responde a criterios funcionales, priorizando aquellos espacios relevantes para los casos de uso definidos en el proyecto, como terapia, telepresencia y teleformación. Asimismo, se ha procurado analizar la cobertura global en las distintas zonas de la Fundación, tanto en exterior como en interior, cubriendo los diferentes módulos y pisos del edificio, y llegando hasta el centro de atención temprana de la Fundación que está justo enfrente de la sede principal (marcado con 16).



Fig. 2.1. Localización de las pruebas de cobertura en la Fundación Juan XXIII

En cada ubicación se realizaron pruebas de conectividad utilizando un módem 5G comercial conectado a la red desplegada. Para cada punto se midieron parámetros de rendimiento en enlace descendente (DL) y ascendente (UL), así como indicadores de calidad de señal.

La Tabla 1 resume los resultados obtenidos en las distintas localizaciones, incluyendo el nivel de cobertura estimado, la distancia aproximada a la antena y los valores de throughput medidos en UL y DL.

Tabla 1. Resumen de resultados de las pruebas de cobertura

ID	Nivel	D(x)	Localización	Zona	DL	UL
1	3	50	Interior	Atención psicológica	145	28.7
2	2	50	Interior	Taller de artesanía	152	28.7
3	1	50	Interior	Gimnasio	180	28.1
4	1	40	Exterior	Patio	143	28.7
5	1	60	Interior	Aula de informática	135	19.5
6	1	10	Interior	Terapia	159	28.7
7	0	15	Exterior	Terraza cafetería	165	28.6
8	0	20	Interior	Piso de entrenamiento	159	28.7
9	0	20	Interior	Auditorio	103	19.5
10	-1	10	Interior	Almacén	64	9.7
11	-1	45	Interior	Taller ocupacional	68	9.6
12	-2	35	Interior	Aulas	-	-
13	-2	70	Exterior	Acceso parking	106	20.7
14	-1	30	Exterior	Playa de camiones	116	21.6
15	0	75	Exterior	Parking visitantes	106	24.3
16	0	110	Exterior	Centro de atención temprana	118	20.9

Los resultados muestran que, en la mayoría de las zonas interiores y exteriores evaluadas, la red proporciona conectividad suficiente para soportar aplicaciones con requisitos de transmisión exigentes, especialmente en enlace ascendente, que es crítico para varios de los casos de uso del proyecto. En aquellas zonas con niveles de señal más reducidos, el rendimiento sigue siendo compatible con servicios de menor demanda de ancho de banda o con tráfico no crítico.

Durante las pruebas se monitorizó de manera continua el tráfico generado y recibido en la red, con el fin de verificar la estabilidad de las mediciones y descartar efectos transitorios. Esta monitorización permite confirmar la consistencia de los valores obtenidos en las

distintas ubicaciones. La siguiente figura muestra el resultado de las pruebas de uplink y downlink para las 16 ubicaciones seleccionadas.



Fig. 2.2. Monitorización del tráfico de red durante las pruebas de cobertura

2.3 Pruebas de slicing

El objetivo de este apartado es verificar el correcto funcionamiento del mecanismo de *network slicing* configurado en la red 5G desplegada, así como la posibilidad de priorizar y limitar dinámicamente la capacidad asignada a un slice concreto. En particular, se evalúa el efecto de la modificación dinámica del porcentaje máximo de recursos asignados a un slice sobre la capacidad efectiva observada en la red.

En estas pruebas se trabaja con un único equipo de usuario (UE), sin tráfico competitivo de otros slices, con el fin de aislar el comportamiento del mecanismo de control de capacidad. El tráfico se genera mediante iperf3 en modo TCP desde la localización 1 (Sala de Atención Psicológica). El UE permanece asociado al mismo slice durante toda la prueba; sin embargo, la configuración de dicho slice se modifica dinámicamente para limitar su capacidad máxima. Las pruebas de enlace ascendente (uplink) y descendente (downlink) se realizan de manera independiente.

2.3.1 Configuración dinámica de slicing en uplink

En esta prueba se evalúa el comportamiento del slicing en enlace ascendente mediante la modificación dinámica del porcentaje máximo de recursos asignados al slice activo. Mientras

el UE genera tráfico TCP de forma continua, se reduce progresivamente la capacidad máxima configurada para el slice, y posteriormente se restablece al valor inicial.

La Figura 2.3 muestra la monitorización del tráfico de uplink durante los cambios de configuración, donde se observa la variación de la capacidad efectiva conforme se modifica el límite máximo del slice.



Fig. 2.3. Monitorización del tráfico de uplink con cambio de configuración de slices

La siguiente table recoge los valores configurados y los valores medidos de throughput para cada uno de los porcentajes aplicados. El orden de las pruebas corresponde a una reducción progresiva de la capacidad máxima del slice, seguida de una restauración final al 100%.

Tabla 2. Medidas de tráfico de uplink para distintas configuraciones de slice

Configurado	Medido	Medido
100%	28.7 Mbps	100%
75%	21.3 Mbps	74.2%
50%	14.4 Mbps	50.2%
25%	7.4 Mbps	25.8%
10%	3.1 Mbps	10.8%
100%	28.7 Mbps	100%

2.3.2 Configuración dinámica de slicing en downlink

De forma análoga a las pruebas de uplink, se realizaron pruebas de control dinámico de slicing en enlace descendente. En este caso, el UE genera tráfico TCP de downlink mientras se modifica en caliente el porcentaje máximo de recursos asignados al slice activo.

La Figura 2.4 presenta la monitorización del tráfico de downlink durante los cambios de configuración del slice, mostrando la adaptación de la capacidad efectiva a los valores configurados.



Fig. 2.4. Monitorización del tráfico de downlink con cambio de configuración de slices

Los resultados cuantitativos obtenidos para cada configuración se resumen en la siguiente tabla, donde se comparan los valores de throughput configurados y medidos para los distintos porcentajes de capacidad máxima asignada.

Tabla 3. Medidas de tráfico de downlink para distintas configuraciones de slice

Configurado	Medido	Medido
100%	168 Mbps	100%
75%	126 Mbps	75.0%
50%	86 Mbps	51.2%
25%	41 Mbps	24.4%
10%	16 Mbps	9.5%
100%	167 Mbps	99.4%

2.3.3 Pruebas en dos módems

Con el objetivo de evaluar el comportamiento del slicing en un escenario con tráfico simultáneo, se realizaron pruebas adicionales utilizando dos equipos de usuario conectados a la red de manera concurrente. En este escenario, cada UE se asocia a un slice distinto, permitiendo observar el reparto de capacidad en función de la configuración de los límites de cada slice. Además del módem en el slice Gold en la ubicación 1, se añade un segundo módem en el slice Bronze en la ubicación 5 (aula de informática).

Durante la prueba, ambos UEs generan tráfico de uplink de forma simultánea, y se monitoriza el tráfico recibido en la red para cada uno de ellos. La Figura 2.5 muestra la evolución temporal del tráfico de uplink para ambos equipos bajo esta configuración.



Fig. 2.5. Monitorización del tráfico de uplink para la configuración con dos módems

La Tabla 4 resume los valores de throughput medidos para cada UE, incluyendo los valores mínimos, máximos y medios observados durante la prueba, así como su correspondencia con los límites configurados en cada slice. Como se puede observar, ambos slices superan la capacidad mínima garantizada (columna “Min”) y se reparten los recursos disponibles de forma proporcional sin llegar al máximo.

Tabla 4. Medidas de tráfico de uplink para la configuración con dos módems

ID	Slice	Pico	Min	Max	Medido	Medido
1	Gold	28.7 Mbps	50% (14 Mbps)	100% (28 Mbps)	17 Mbps	86%
5	Bronze	19.5 Mbps	5% (1 Mbps)	20% (4 Mbps)	2.7 Mbps	14%

Pruebas de estabilidad

El objetivo de las pruebas de estabilidad es evaluar el comportamiento de la red 5G desplegada bajo carga periódica y sostenida en el tiempo, verificando que su funcionamiento se mantiene estable durante un periodo prolongado de operación. Para ello, se configuró un conjunto de pruebas automáticas de generación de tráfico que se ejecutan de forma recurrente a lo largo del día.

Las pruebas se realizaron durante aproximadamente un mes, desde el 28 de noviembre hasta el 20 de diciembre, ejecutándose de manera periódica cada hora, 24 horas al día. En cada ejecución se generan flujos de tráfico de máxima carga mediante iperf3, con una duración de tres minutos por prueba. Se utilizan dos equipos de usuario: uno ubicado en la posición 1, asociado al slice Gold, y otro en la posición 5, asociado al slice Bronze.

La Figura 2.6 muestra un ejemplo del comportamiento observado durante una de las ventanas horarias de prueba, incluyendo el tráfico correspondiente a ambos slices. En esta figura, *slicetun1* corresponde al tráfico del slice Gold y *slicetun3* al tráfico del slice Bronze.



Fig. 2.6. Detalle de la configuración de carga en las pruebas de estabilidad

A continuación, se presentan las métricas monitorizadas durante todo el periodo de observación. La Figura 2.7 muestra la evolución temporal del tráfico generado en la red a lo largo de las distintas ejecuciones de las pruebas. Se observa un patrón repetitivo asociado a la ejecución periódica de las pruebas de carga.

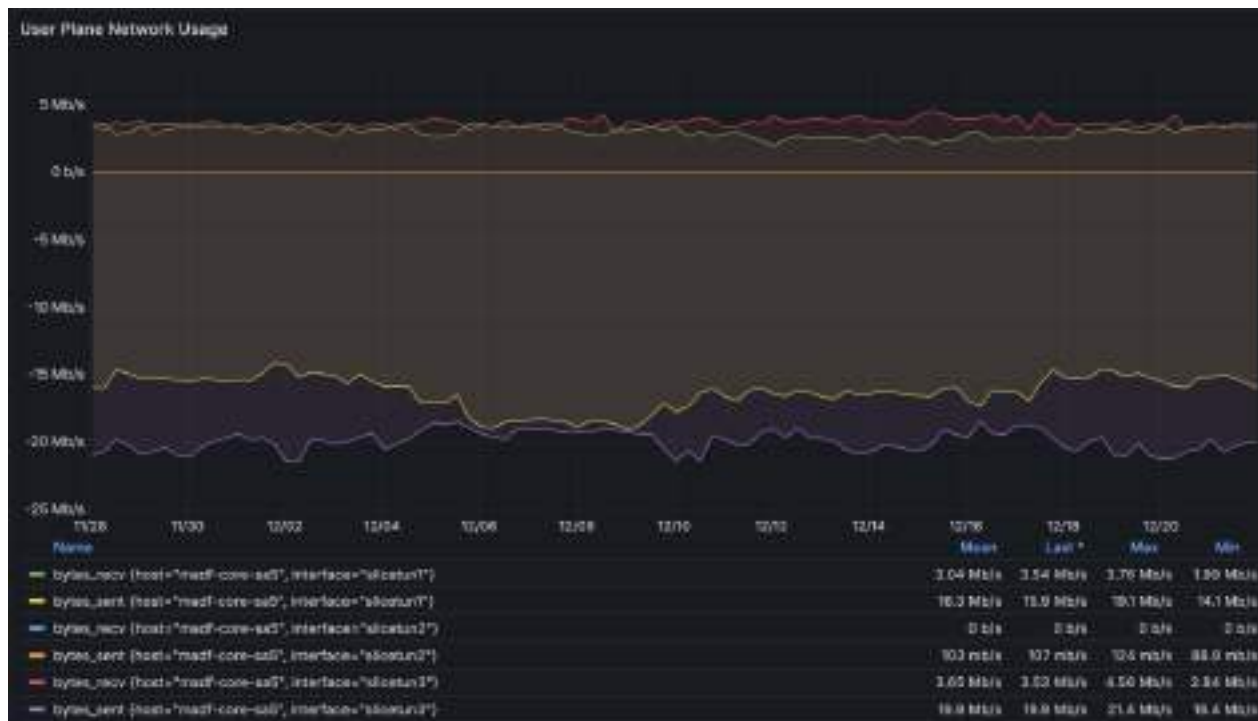


Fig. 2.7. Monitorización del tráfico en el core (UL/DL, dos slices) durante las pruebas de estabilidad

La latencia se monitoriza mediante medidas de RTT desde el módem ubicado en la posición 1 hasta el core de la red. La evolución temporal de esta métrica durante todo el periodo se muestra en la Figura 2.8. El mínimo (23 ms) corresponde a los momentos en los que no hay carga de tráfico y el máximo (561 ms) a los momentos cargados. Más Adelante se analizará el detalle.

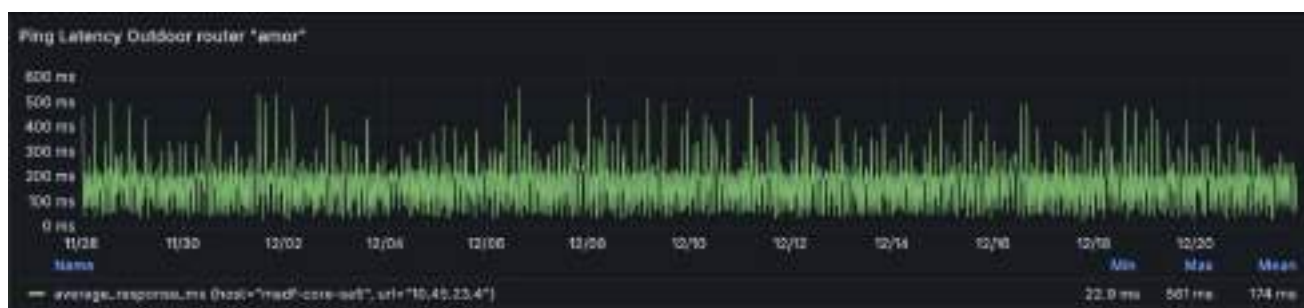


Fig. 2.8. Monitorización del tiempo de ping (RTT) desde un módem durante las pruebas de estabilidad

La Figura 2.9 presenta la carga de CPU de la máquina virtual que ejecuta el core 5G basado en open5gs. Esta métrica permite observar el impacto de la generación periódica de tráfico sobre el plano de control y de usuario del núcleo de red. Se observa que el consume de CPU es bastante estable.



Fig. 2.9. Monitorización de la carga de CPU del core durante las pruebas de estabilidad

El consumo energético de la estación base radio se muestra en la Figura 2.10. En esta figura se distinguen las contribuciones correspondientes a la banda base y al elemento de radio remoto (RRH), permitiendo analizar su comportamiento durante los periodos con y sin tráfico activo.

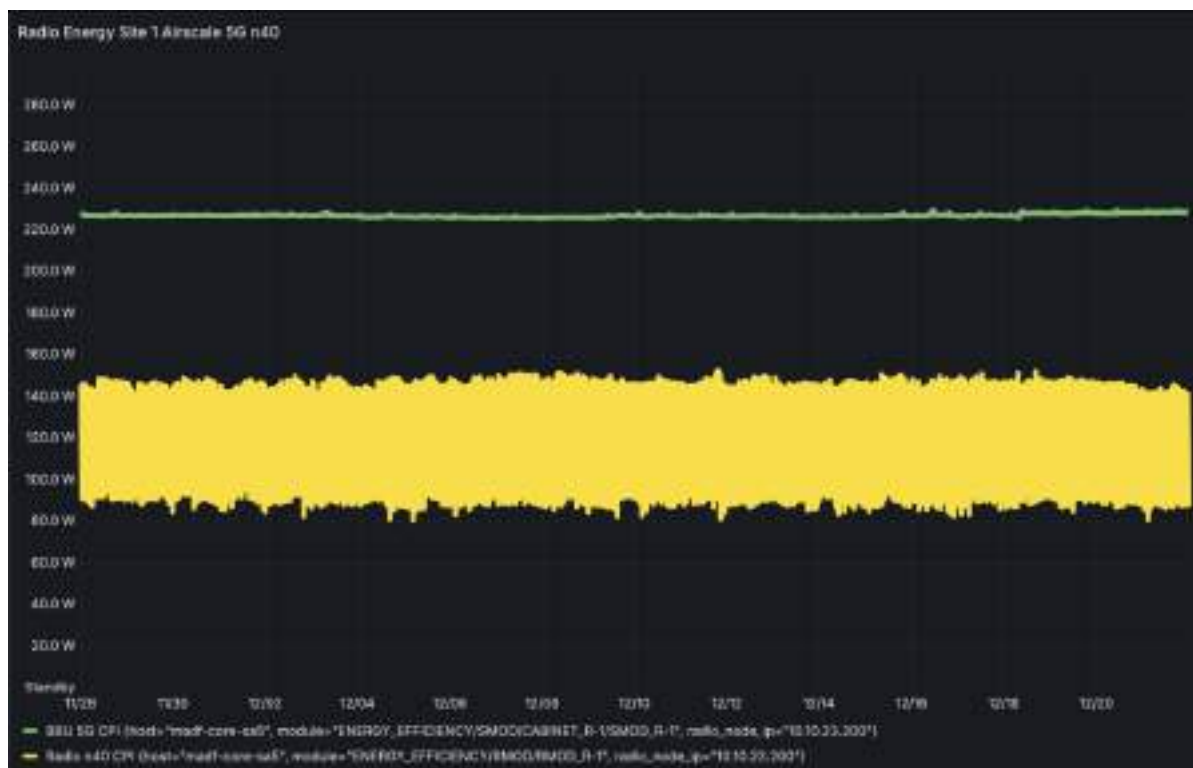


Fig. 2.10. Monitorización del consumo de energía en la banda base (verde) y RRH (amarillo) durante las pruebas de estabilidad

La calidad de la señal radio se monitoriza en paralelo para los dos equipos de usuario utilizados en la prueba. Las Figuras 2.11 y 2.12 muestran la evolución temporal de los valores de RSRP, RSRQ y SINR medidos en las posiciones 1 y 5.

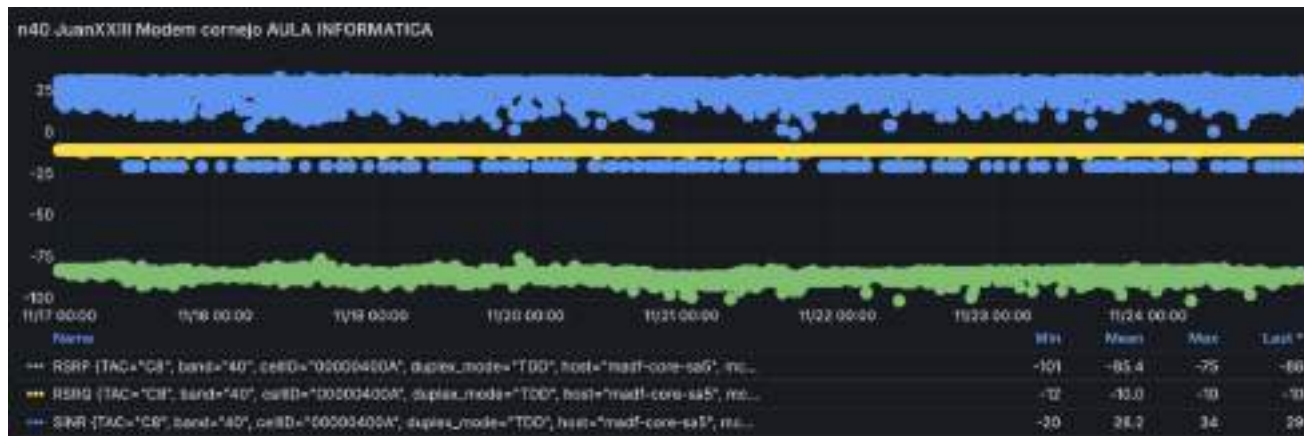


Fig. 2.11. Monitorización de la Calidad de señal del módem en posición 5 (aula de informática) durante las pruebas de estabilidad

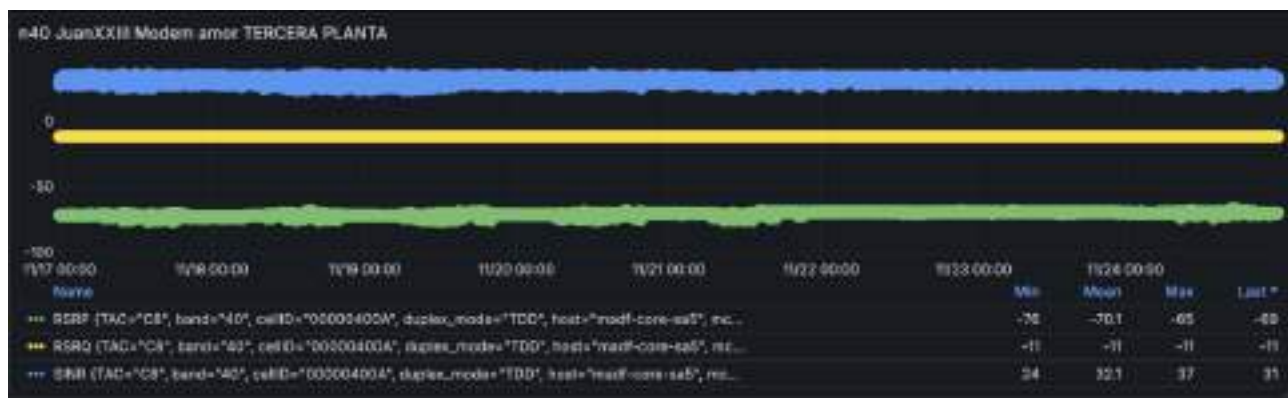


Fig. 2.12. Monitorización de la calidad de señal del módem en posición 1 (despachos tercera planta) durante las pruebas de estabilidad

Para facilitar un análisis más detallado y correlacionado de las distintas métricas, las Figuras 2.13 a 2.17 presentan un zoom temporal sincronizado de cuatro horas, comprendido entre las 23:00 y las 03:00. En este intervalo se muestran, de forma alineada en el tiempo, el tráfico generado, la latencia, la carga de CPU del core, el consumo energético de la radio y la calidad de señal radio.



Fig. 2.13. Detalle del tráfico en el core (UL/DL, dos slices) durante las pruebas de estabilidad

En particular, se observa que el RTT se ve afectado de forma muy notable en situaciones de carga, pudiendo tener picos de hasta un segundo si debe competir con el tráfico de iperf3. Es conveniente destacar, no obstante, que iperf3 satura la capacidad de la red; en una situación de alta carga pero sin saturación, el RTT se incrementaría en menor medida.



Fig. 2.14. Detalle del tiempo de ping (RTT) desde un módem durante las pruebas de estabilidad

Las medidas de calidad de señal en los módems muestran que la SINR oscila fluctúa a lo largo del tiempo, incluso sin mover el módem ni cambiar la potencia recibida.



Fig. 2.15. Detalle de la calidad de señal del módem en posición 5 (aula de informática) durante las pruebas de estabilidad



Fig. 2.16. Detalle de la calidad de señal del módem en posición 1 (despachos tercera planta) durante las pruebas de estabilidad

Finalmente, la medida de energía muestra el incremento en la potencia de la cabecera radio remota (RRH) cuanto transmite o, en menor medida, cuando recibe tráfico de los módems con respecto a su estado de reposo. La banda base, por el contrario, mantiene un consumo homogéneo a lo largo del tiempo.



Fig. 2.17. Detalle del consumo de energía en banda base (verde) y RRH (amarillo) durante las pruebas de estabilidad

2.5

Conclusiones

Las pruebas presentadas en este capítulo permiten validar la adecuación de la red 5G desplegada en banda n40 como infraestructura de comunicaciones para el proyecto INCLUVERSO 5G. En particular, la red proporciona cobertura en las ubicaciones relevantes de la Fundación Juan XXIII Roncalli en las que se desarrollan los casos de uso del proyecto.

Los ensayos realizados sobre el mecanismo de *network slicing* confirman su correcto funcionamiento y la posibilidad de modificar dinámicamente los parámetros de configuración asociados a un slice. Los cambios aplicados en caliente sobre la capacidad máxima asignada se reflejan de forma directa en la capacidad efectiva observada, tanto en enlace ascendente como en enlace descendente.

Asimismo, en escenarios con múltiples equipos de usuario conectados simultáneamente, el reparto de capacidad entre slices se ajusta a la configuración establecida, permitiendo diferenciar el comportamiento de los distintos flujos de tráfico.



En conjunto, los resultados obtenidos en este capítulo permiten considerar la red 5G desplegada como una infraestructura adecuada para dar soporte a la validación de los casos de uso del proyecto y a los análisis posteriores de calidad de servicio y calidad de experiencia.

3 Validación de los modelos y técnicas para modelar y controlar QoS/QoE

Modelos paramétricos de QoE para los escenarios de aplicación del proyecto

Los modelos paramétricos se emplean para planificar las capacidades de un sistema de comunicaciones y poder estimar la QoE esperada a partir de las medidas y estadísticas de tráfico y computación del sistema.

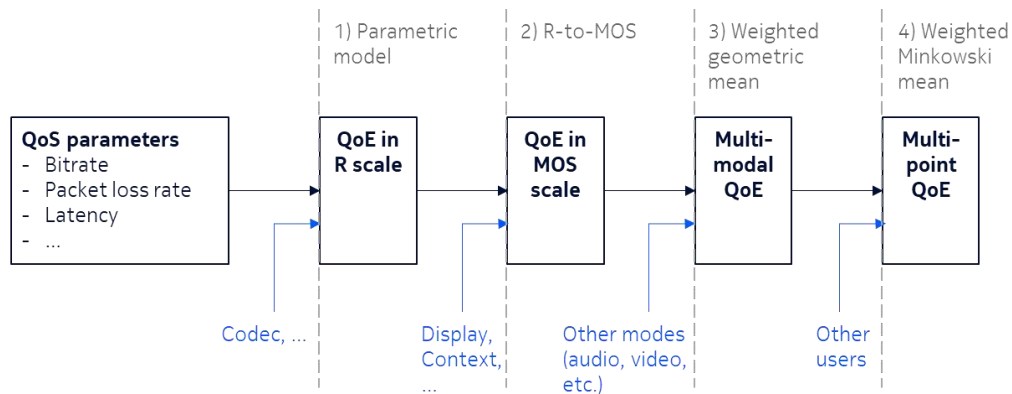


Fig. 3.1. Descripción general de la lógica del modelo modular de QoE

$$MOS(R, \sigma) = 5 - \sum_j^4 E\left(\frac{j + 0.5 - \mu(R)}{\sigma}\right) \quad (3.1)$$

donde $E()$ es la función de distribución normal:

$$E(x) = \frac{1}{\sqrt{2\pi}} \int_{-\infty}^x \exp\left(-\frac{t^2}{2}\right) dt \quad (3.2)$$

Y $\mu(R)$ representa la relación entre R y la calidad subjetiva en una escala continua. Esta ecuación permite pasar de la escala ilimitada de μ a la escala acotada de MOS, como se indica en la siguiente figura.

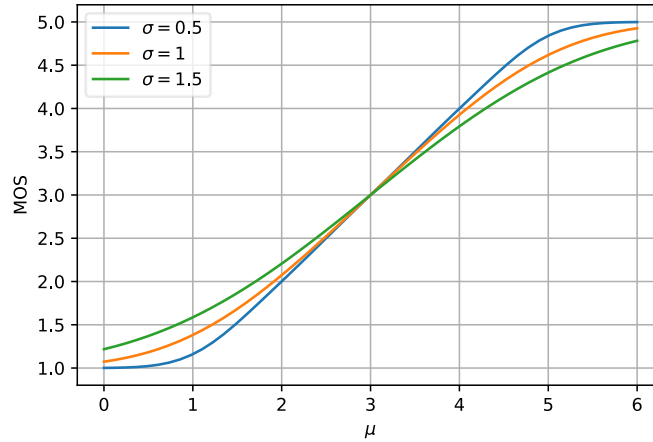


Fig. 3.2. Relación entre μ y MOS para diferentes valores de σ

3.1.1 Escenario 1: Teleportación Virtual (Virtual Teleportation)

El principal componente de la QoE en este escenario es el vídeo inmersivo transmitido por el búho. Por ello, el modelo de calidad se centrará en este flujo. Hay dos principales KPIs que determinan la calidad del vídeo recibido: el bitrate de codificación y las pérdidas de paquetes.

Para el bitrate de codificación, usaremos la ecuación (3.1) para determinar $MOS_{BR}(R, \sigma)$. Fijamos $\sigma = 0.7$ (adecuado para video inmersivo) y $R = \log(BR)$, como se determine en E2.2. A partir de las pruebas subjetivas (ver E3.3) fijamos los umbrales de percepción y aceptación en 6 y 3 Mbps respectivamente, obteniendo el valor de $\mu(R)$:

$$\mu(BR) = 2.164(\log(BR(kbps)) - \log(3000)) + 3 \quad (3.3)$$

Para la tasa de pérdida de paquetes, se usa $R = \sqrt{PLR(\%)}$, como se propone en Díaz et al. (2020). Se usa el mismo valor de $\sigma = 0.7$ y se estima una tasa de aceptación del 3%, haciendo converger la tasa del 0% a un MOS=5. Con ello queda:

$$\mu(PLR) = -1.766(\sqrt{PLR(\%)} - \sqrt{3}) + 3 \quad (3.4)$$

Dado que el efecto de las pérdidas se impone sobre la mala calidad del vídeo para valores relevantes de la PLR, se toma el mínimo de ambos para determinar el MOS. Esto es el equivalente a una media de Minkowski con $p = -\infty$.

$$MOS = \min(MOS_{BR}, MOS_{PLR}) \quad (3.5)$$

La Fig. 3.3 muestra las curvas de MOS para el bitrate y la tasa de pérdidas.

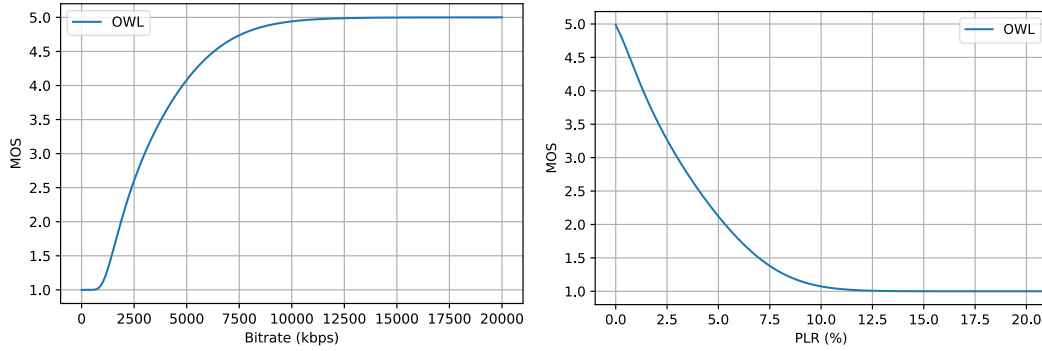


Fig. 3.3. Curvas de MOS para el escenario de teleportación virtual

3.1.2 Escenario 2: Procesamiento delegado de Realidad Mixta (*Mixed Reality computation offloading*)

Para el análisis de este escenario, nos centraremos en el algoritmo de segmentación semántica de Thundernet. Para un análisis más detallado de los distintos algoritmos y sus propiedades se puede consultar el capítulo 4 de este documento (para analizar el rendimiento en el MEC), o el entregable E3.3 (donde se describen sus resultados a nivel de experiencia de usuario).

El principal componente de la QoE en este escenario es el vídeo segmentado en el servidor del edge. Por ello, el modelo de calidad se centrará en este flujo. Hay dos principales KPIs que determinan la calidad del vídeo recibido: el tiempo total de ida y vuelta (RTT), considerando uplink del frame capturado, segmentación y downlink de la máscara de segmentación; y la resolución del frame con el que se trabaja.

Para calcular MOS_{RTT} se usa el modelo propuesto por Cortés et al. (2023), que se puede ajustar a un modelo lineal en la escala R, como se describe en el E2.2. A partir de las pruebas subjetivas realizadas, determinamos unos umbrales de percepción y aceptación de 45 y 75 ms respectivamente, lo que da lugar al siguiente modelo, con $\sigma = 0.8$:

$$\mu(RTT) = -0.05(RTT(\text{ms}) - 75) + 3 \quad (3.6)$$

Los valores para la resolución se obtienen a partir del análisis de las pruebas subjetivas, con $\sigma = 0.8$ y $\mu(R)$:

$$\mu(RES) = 1.154(\log(W \times H) - 12.635) + 3.7 \quad (3.7)$$

Esto ajusta aproximadamente la calidad a un valor de 3.7 para una resolución de 640x480 y cercano a 4.5 para la resolución de trabajo (ver E3.3) de 1280x480. Se combinan ambos valores mediante la media armónica ($p = -1$):

$$MOS = (MOS_{RTT} \times MOS_{RES})^{1/2} \quad (3.8)$$

Validamos estas decisiones viendo que, en las condiciones de ejecución del caso de uso de terapia (vídeo de 1280x480, RTT en torno a 50 ms), se obtiene un MOS = 4.3, similar al reportado por los usuarios (ver E3.3).

La Fig. 3.4 muestra las curvas de MOS para los dos KPIs bajo análisis.

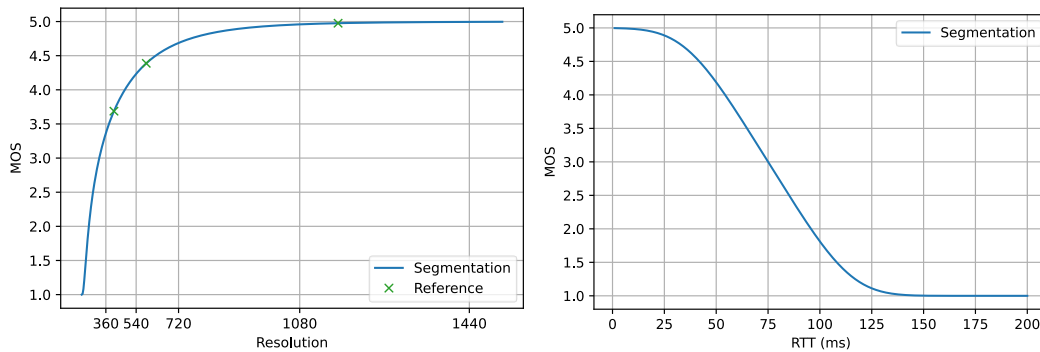


Fig. 3.4. Curvas de MOS para el escenario de procesamiento delegado

Para poder analizar el efecto del RTT en la QoE es necesario considerar, además de la transmisión de los paquetes, el tiempo de procesamiento en el servidor. Como se muestra en el capítulo 4, este tiempo depende del algoritmo y de la resolución. Para poder analizar la QoE de la red según el modelo, es preciso tener un modelo estocástico de estos tiempos. Emplearemos un modelo sencillo y lineal, basado en los tiempos analizados para Thundernet, como se muestra en la Fig. 3.5. Para el uso de otro tipo de algoritmos, este modelo debería adaptarse a cada uno de ellos.

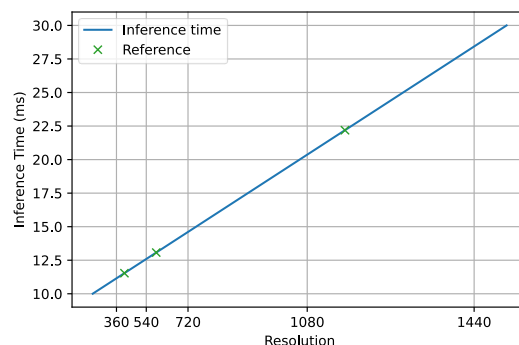


Fig. 3.5. Tiempo de inferencia de referencia para cada resolución

3.1.3 Escenario 3: Regulación Emocional (*Emotional Regulation*)

Para el escenario de regulación emocional, tanto la transmisión de datos como los requisitos de latencia son mucho menores. En todo caso, el principal KPI que se debe medir es la latencia de ida y vuelta. En el proyecto no se dispone de suficientes datos que permitan alimentar un modelo así que, de acuerdo al diseño propuesto en E2.2, el modelo se diseñará a partir del

conocimiento del estado del arte. Dado que la incertidumbre del modelo (y, potencialmente, de los valores de MOS) es mayor, se asigna un valor relativamente alto de $\sigma = 1.0$. Como ya se determinó en E2.2, para la QoE de la respuesta en el tiempo es aproximadamente lineal en escala R. Según muestra la literature científica en sistemas de regulación emocional en tiempo real (Braithwaite et al., 2013; Benedek & Kaernbach, 2010), se puede establecer el umbral perceptual en unos 300 ms y el umbral de aceptabilidad en 2s. Con esto se obtiene:

$$\mu(RTT) = -0.001(RTT(\text{ms}) - 2000) + 4 \quad (3.9)$$

Que se representa en la figura:

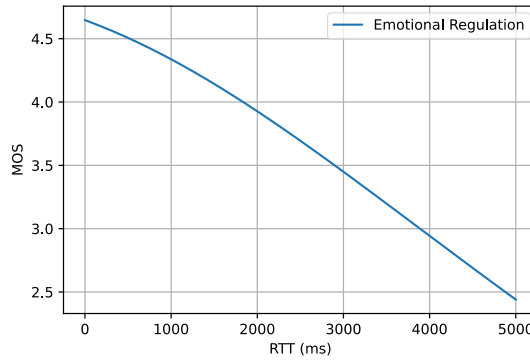


Fig. 3.6. Curva de MOS para el escenario de regulación emocional

3.2 Validación de las mejoras en el emulador FikoRE

Como se ha mencionado en apartados anteriores, el principal desarrollo en el emulador FikoRE consistió en un análisis profundo y refinamiento de la capa física del emulador FikoRE, con un enfoque particular en mejorar el realismo y la consistencia interna de la estimación de señal. Esto permite que las aplicaciones de las herramientas de QoS del emulador, como el establecimiento de métricas de capa MAC (scheduling) y de priorización (pesos), proporcionen resultados significativos.

El proceso comenzó con una revisión exhaustiva de la versión anterior del emulador (González-Morín et al., 2023), durante la cual se identificaron varias limitaciones. Uno de los problemas clave fue la falta de consistencia entre los cálculos de enlace ascendente (UL) y enlace descendente (DL): era posible que el mismo equipo de usuario (UE) experimentara LOS en el enlace ascendente y NLOS en el enlace descendente, ya que los dos enlaces se calculaban independientemente.

Para resolver esto, se implementó una nueva clase PHY compartida (phy_shared), asegurando clasificación coherente LOS/NLOS, pérdidas de penetración y sombreado para tanto UL como DL. Este cambio también permitió un modelado más realista de la asimetría de enlace considerando diferencias en potencia TX, ganancia de antena, figura de ruido e interferencia experimentada en el UE y la estación base (BS). Como resultado, los enlaces UL

y DL ahora se comportan de manera más consistente y realista, como se espera en un despliegue del mundo real.

Otro problema identificado fue la ausencia de correlación espacial en el desvanecimiento por sombra. Anteriormente, dos UEs ubicados en las mismas coordenadas podían experimentar efectos de sombreado diferentes. Esto se abordó generando mapas de desvanecimiento por sombra 2D, habilitando sombreado espacialmente correlacionado y asegurando que usuarios cercanos experimenten condiciones de canal similares.

Más allá de estas correcciones fundamentales, se incorporaron varias mejoras en la capa PHY para aumentar la fidelidad del modelado:

- Diferenciación de ganancia de antena y figura de ruido entre nodos
- Potencia de transmisión UL adaptativa a la distancia
- Modelado de interferencia de usuarios concurrentes
- División de potencia a través de capas espaciales MIMO
- Modelos de desvanecimiento multitrayecto
- Efectos de absorción de oxígeno en la banda de 52–67 GHz
- Modelado de pérdidas de penetración exterior-interior

En la capa MAC, no se realizaron modificaciones estructurales. Sin embargo, el funcionamiento del planificador se probó exhaustivamente bajo la nueva lógica de estimación SINR. También se automatizó todo el proceso de lanzamiento del emulador y ejecución de scripts de análisis para post-procesamiento de registros. Además, se revisó el cálculo de eficiencia RF, asegurando consistencia con los estándares 3GPP. Esto incluyó corregir suposiciones sobre el ancho de banda total utilizable en la BS, reconociendo que no todo el ancho de banda está disponible para transmisión debido a bandas de guarda, canales de control y limitaciones de recursos.

Para validar estas mejoras, se llevó a cabo una campaña de medición 5G integral a través de múltiples escenarios del mundo real. Estos incluyeron entornos interiores, microceldas urbanas y macroceldas rurales, con mediciones tomadas tanto al aire libre como dentro de vehículos. Los experimentos se diseñaron cuidadosamente para exponer FikoRE a condiciones de propagación realistas como transiciones LOS/NLOS, variación de altura del usuario y pérdidas de penetración. Las trazas recopiladas incluyeron métricas como SINR, RSRP, throughput, potencia TX y MCS a lo largo del tiempo. Se describen ahora los resultados más relevantes de esa campaña.

La validación se enfocó en dos aspectos clave. Primero, verificar el comportamiento del modelo de pérdidas de trayecto ABG implementado contra datos del mundo real. La figura siguiente muestra el exponente de pérdidas de trayecto α , extraído de mediciones RSRP en

una microcelda urbana. El α estimado se alinea bien con los valores utilizados en el emulador, demostrando la validez del modelo.

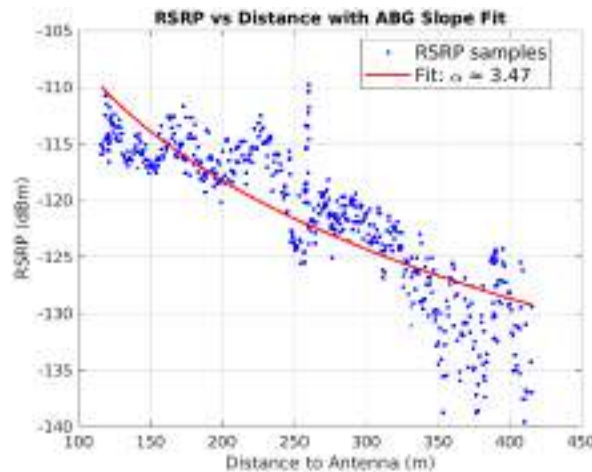


Fig. 3.7. Parámetro de desvanecimiento en microcelda urbana (UMi): modelo frente a medidas

Segundo, evaluamos la capacidad del emulador para replicar la variabilidad estadística del desvanecimiento por sombra. En un escenario rural NLOS, derivamos la desviación estándar de sombreado de registros RSRP, aplicando normalización basada en ABG a una distancia de referencia de 730,93 metros. La desviación estándar estimada fue 6,24 dB, que coincide estrechamente con los 6,7 dB utilizados en la configuración rural NLOS del emulador.

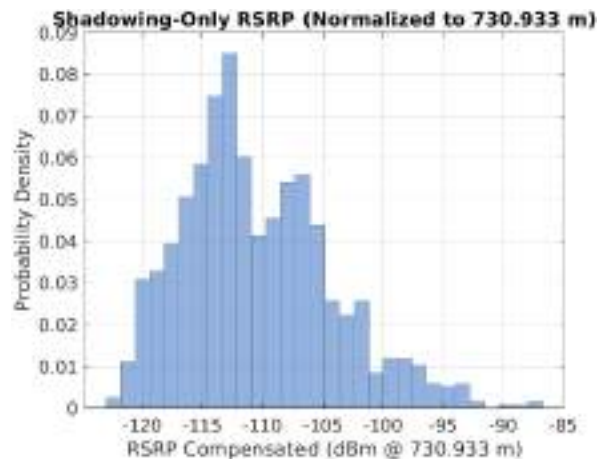


Fig. 3.8. Distribución de los valores de sombreado obtenidos en una campaña de macroceldas rural (RMa)

Finalmente, para evaluar la mejora general en la estimación de señal, se compararon las distribuciones SINR de tres fuentes: la versión anterior del emulador, la versión mejorada y mediciones del mundo real. Este análisis comparativo demuestra cómo el nuevo enfoque de

modelado—especialmente el uso de sombreado correlacionado y pérdidas de trayecto refinadas—acerca significativamente FikoRE a la realidad.

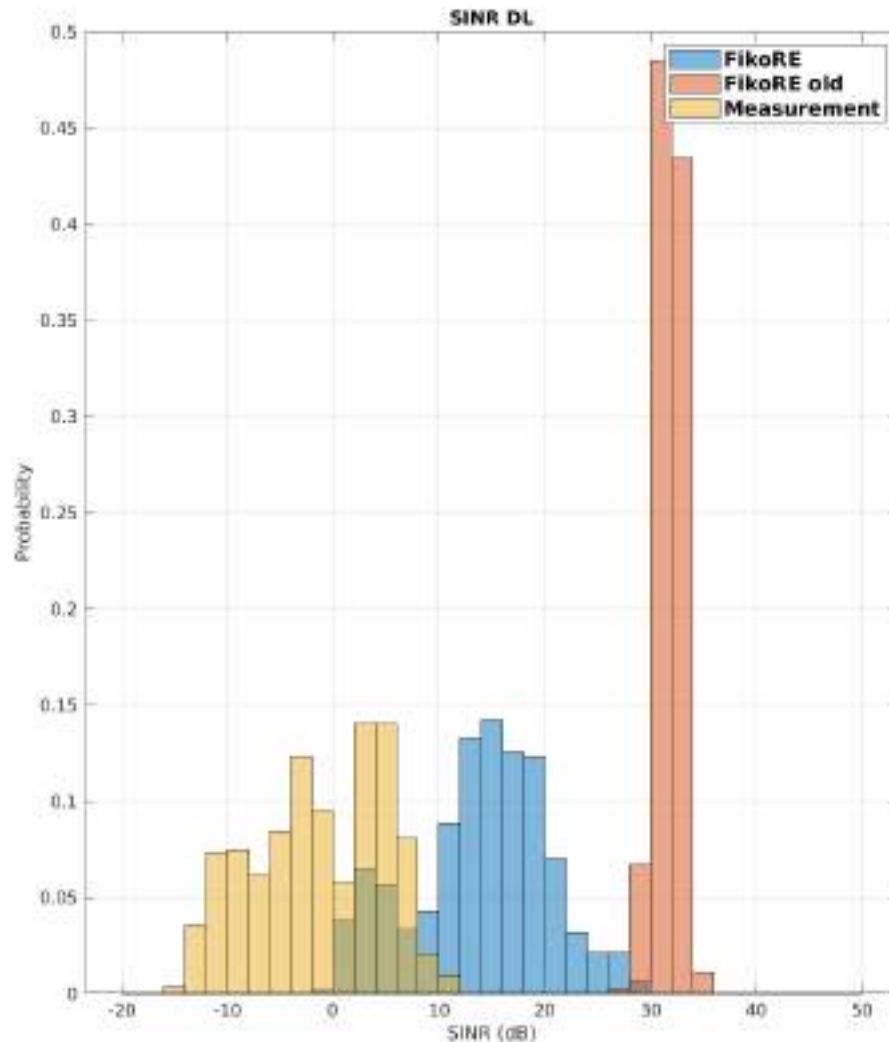


Fig. 3.9. Distribución de SINR: medida y estimada por ambas versiones de FikoRE

Esta figura ilustra la alineación mejorada del emulador con mediciones de campo en un escenario de microcelda urbana densa bajo condiciones NLOS a 2,38 GHz.

Adicionalmente, el experimento de macrocelda rural (RMa NLOS a 3,5 GHz), realizado como una caminata aleatoria a través de un pueblo, proporciona otra perspectiva sobre el comportamiento del emulador. Aquí, observamos cómo el throughput de enlace descendente predicho por la nueva versión de FikoRE comienza a seguir estrechamente el de la red real.

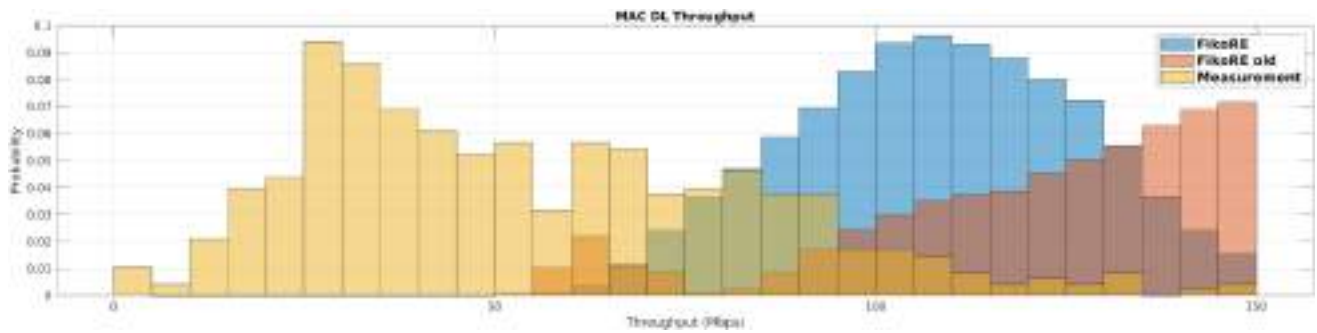


Fig. 3.10. Distribución de tráfico alcanzado en entorno rural: real vs. emulado

Es importante notar que estos resultados son específicos del experimento. En cada situación, las condiciones de señal son altamente variables y no pueden ser completamente capturadas por un escenario genérico. El modelado preciso de un despliegue del mundo real requeriría un entorno 3D completo, incluyendo propiedades de materiales, trayectos de reflexión, fuentes de interferencia y densidad de usuarios activos. Afortunadamente, este nivel de detalle no es el objetivo de FikoRE.

En su lugar, FikoRE está diseñado para servir como un emulador a nivel de sistema, ideal para probar el comportamiento de aplicaciones, protocolos y estrategias de planificación bajo condiciones radio realistas pero abstraídas. Las mejoras introducidas durante este trabajo acercan el emulador un paso más a cerrar la brecha entre simulación y validación de campo—proporcionando tanto confiabilidad como usabilidad para futura investigación y desarrollo de aplicaciones.

3.3

Pruebas en red emulada (5QIs)

El objetivo de este apartado es analizar la viabilidad de los escenarios de aplicación del proyecto sobre red 5G cuando se introducen distintos niveles de carga y mecanismos de gestión de calidad de servicio, comparando los resultados obtenidos con los requisitos definidos en el entregable E1.2. Para ello, se emplea el emulador FikoRE, que permite reproducir distintos tipos de celda, bandas de frecuencia y números de usuarios concurrentes de forma controlada.

Las pruebas se estructuran según los tres escenarios base definidos en E1.2: teleportación virtual, procesamiento delegado de realidad mixta y regulación emocional. Para cada uno de ellos se consideran los escenarios de emulación correspondientes, variando el tipo de celda, la banda de frecuencia y el número de usuarios, así como la prioridad asignada mediante distintos valores de 5QI.

El análisis se centra en métricas de red (throughput, latencia y pérdidas), que se comparan directamente con los requisitos del escenario correspondiente. A partir de estas métricas se estima la QoE mediante los modelos descritos en el apartado 3.1, lo que permite evaluar el

cumplimiento de los objetivos de experiencia definidos en E1.2 sin necesidad de realizar pruebas subjetivas adicionales.

Tabla 5. Escenarios de simulación y requisitos.

	Teleportación Virtual	Computación Delegada en MR	Regulación Emocional
Tipo de celda	Microcelda Urbana	Picocelda urbana / Hostpot	Macrocela Urbana
Banda n40	20 MHz TDD	20 MHz TDD	20 MHz TDD
Banda n78	100 MHz TDD	100 MHz TDD	-
Banda n256	-	800 MHz TDD	-
Distancia a la antena	150 m	50 m	500 m
Usuarios concurrentes	1-10	1-10	10-50
Distancia media de los usuarios	200 m	50 m	1000 m
Distancia máxima de los usuarios	500 m	200 m	5000 m
Target MOS	3.5	4.0	
DL Rate	0.3-0.8 Mbps	1-20 Mbps	0.1 Mbps
UL Rate	5-20 Mbps	10-200 Mbps	0.1 Mbps
DL PDB	100 ms	20 ms	300 ms
UL PDB	100 ms	20 ms	300 ms
DL PER	FFS	FFS	FFS
UL PER	FFS	FFS	FFS
Tipo Tráfico	Constante (por frame)	Constante (por frame)	A ráfagas

En cada escenario de emulación se define un **usuario emulado**, que representa el servicio de interés del caso de uso, y un conjunto de usuarios adicionales que generan tráfico concurrente. El número de usuarios concurrentes y su distribución espacial se ajustan a los valores definidos en los escenarios de simulación de referencia de E1.2.

El tráfico generado por el usuario emulado reproduce las características del servicio correspondiente (tasas binarias, periodicidad y sensibilidad a latencia), mientras que el resto de usuarios generan tráfico de fondo con parámetros homogéneos. La gestión de QoS se implementa mediante la asignación de distintos valores de 5QI al usuario emulado, permitiendo analizar el impacto de la priorización frente a escenarios sin gestión específica.

Las métricas obtenidas se comparan con los requisitos de conectividad definidos en E1.2 para cada escenario, y se traducen posteriormente a valores de MOS estimado utilizando los modelos paramétricos de QoE del proyecto.

3.3.1 Escenario 1: Teleportación Virtual (Virtual Teleportation)

Este escenario corresponde a un servicio con tráfico dominante en enlace ascendente, asociado a la transmisión de vídeo inmersivo en tiempo real. Las pruebas se realizan en un entorno de microcelda urbana, considerando distintas bandas de frecuencia (n40 y n78) y un número variable de usuarios concurrentes.

Tabla 6. Parámetros de simulación del escenario de teleportación virtual en red emulada.

Microcelda Urbana	Escenario E1.2	Simulación	
Banda		n40	n78
Ancho de Banda		20 MHz TDD	100 MHz TDD
Usuario emulado			
Distancia a la antena	150 m	50m	150m
UL Rate	6.17 Mbps		
Tipo Tráfico	Constante (por frame)		
Usuarios adicionales			
Usuarios concurrentes	1-10	1, 5	1, 10
Distancia media de los usuarios	200 m	100m	200m
Distancia máxima de los usuarios	500 m	200m	500m
UL Rate	5-20 Mbps	6 Mbps	6 Mbps

La siguiente figura muestra la trayectoria e histogramas de SINR para una simulación con 10 usuarios. El usuario emulado es UE0, estático en la posición (0, 50) de la celda.

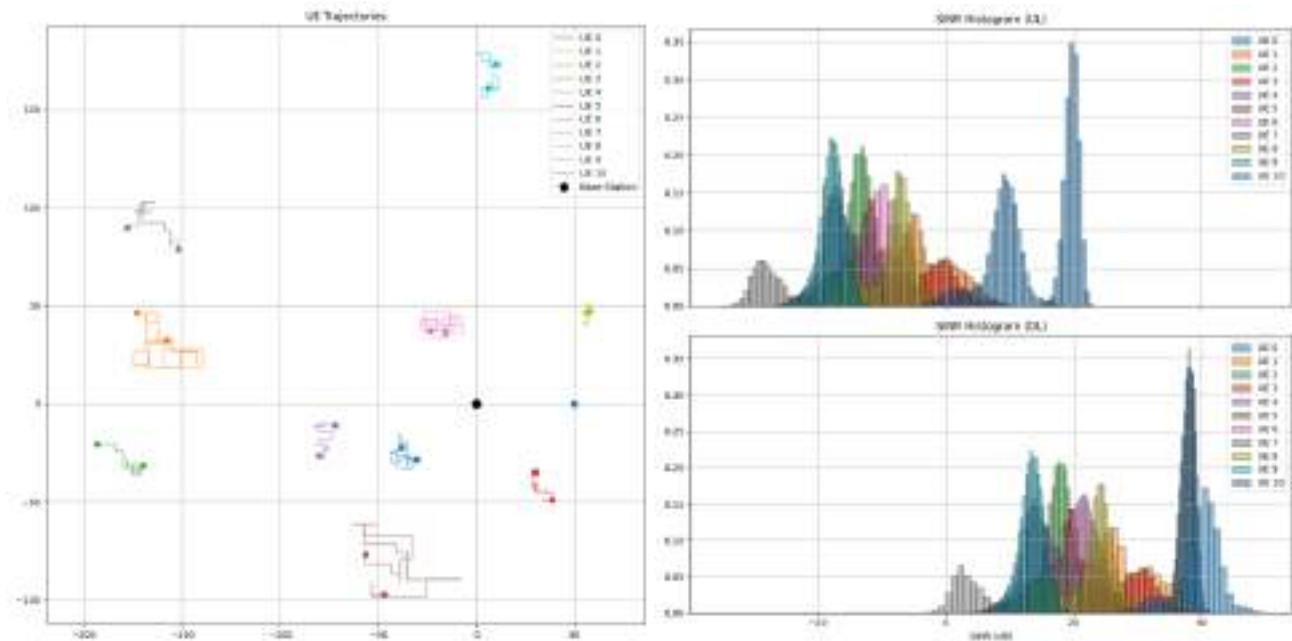


Fig. 3.11. Trayectorias e histogramas de SINR para una simulación con 10 usuarios

Los resultados obtenidos muestran el comportamiento del throughput de uplink y de la latencia de paquete para el usuario emulado a medida que aumenta la carga en la celda. Estos valores se comparan con los requisitos de tasa binaria mínima y presupuesto de retardo definidos en E1.2 para este escenario. Como se muestra en la siguiente figura, cuando solo hay un usuario concurrente, la red es capaz de soportar el tráfico de uplink del *snowl*. Sin embargo, con más usuarios (5 en n40, 10 en n78), aparecen pérdidas muy significativas. La priorización del usuario emulado con un 5QI=6 subsana este error, aunque no garantiza completamente los recursos (hay algunas pérdidas en el escenario en n40).

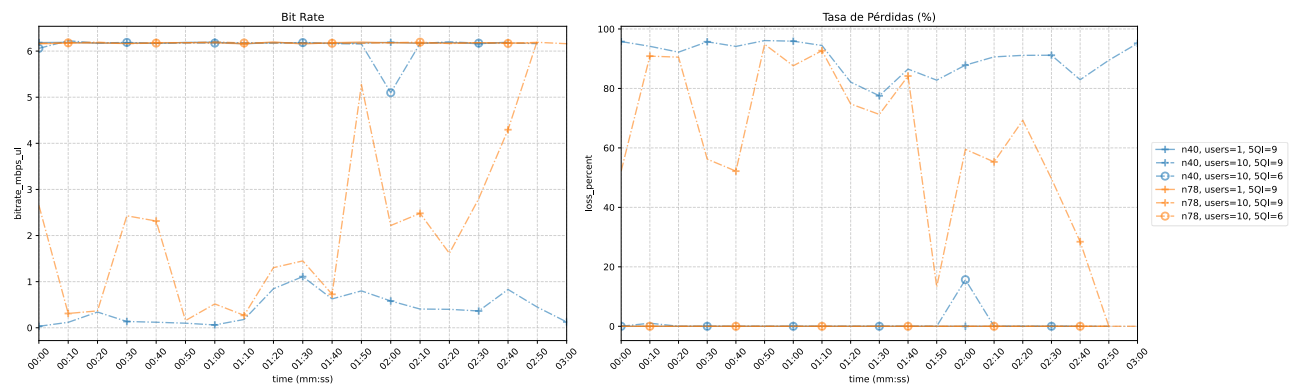


Fig. 3.12. Throughput y latencia de uplink del usuario emulado en función del número de usuarios concurrentes.

A partir de las métricas medidas, se calcula el MOS estimado para el servicio de teleportación virtual. Este valor permite evaluar el impacto conjunto de las limitaciones de red sobre la experiencia percibida, identificando las configuraciones en las que se alcanza o no el objetivo de MOS definido en los requisitos. Como se observa en la table siguiente, el uso de un 5QI priorizado permite cumplir el requisito de MOS > 3.5, manteniendo la latencia dentro de los márgenes especificados.

Tabla 7. MOS estimado para el escenario de teleportación virtual bajo distintas configuraciones de carga y 5QI.

	Objetivo	n40			n78		
Usuarios concurrentes	1-10	1	5	5	1	10	10
5QI		9	9	6	9	9	6
MOS	3.5	4.4	1.0	4.3	4.4	1.0	4.4
UL Bit Rate	6.18	6.18	0.40	6.11	6.18	2.07	6.18
UL Packet Delay	PDB < 100 ms	8.6	340	32.1	7.3	300	7.0
UL Loss Rate (%)	FFS	0.0	90.3	0.9	0.0	62.4	0.0

3.3.2 Escenario 2: Procesamiento delegado de Realidad Mixta (*Mixed Reality computation offloading*)

En este escenario el servicio de interés combina tráfico de subida y bajada, y presenta una elevada sensibilidad a la latencia extremo a extremo, ya que el retardo total condiciona directamente la tasa de cuadros efectiva percibida por el usuario. Las pruebas se realizan en un entorno de picocelda urbana o hotspot, considerando las bandas n40, n78 y n256.

Tabla 8. Parámetros de simulación del escenario de procesamiento delegado

Hostpot	Escenario E1.2	Simulación		
Banda		n40	n78	n256
Ancho de Banda		20 MHz TDD	100 MHz TDD	800 MHz TDD
Usuario emulado				
Distancia a la antena	50 m	50 m	50 m	50 m
DL Rate	1-20 Mbps	10.5 Mbps		
UL Rate	10-200 Mbps	10.5 Mbps		

Tipo Tráfico	Constante (por frame)	Constante (por frame) / 30 FPS		
Usuarios adicionales				
Usuarios concurrentes	1-10	1, 5	1, 10	1, 10
Distancia media de los usuarios	50 m	50m	20m	20m
Distancia máxima de los usuarios	200 m	200m	100m	100m
DL Rate	1-20 Mbps	5 Mbps	10 Mbps	10 Mbps
UL Rate	10-200 Mbps	5 Mbps	10 Mbps	10 Mbps

La siguiente figura muestra las trayectorias de dos pruebas en bandas n40 y n78 respectivamente. Ambas son pruebas simulando un escenario *hotspot*, con usuarios y antena en interiores. Para la prueba en n78 (y en n256) se ha optado por un escenario más denso, con usuarios muy cerca de la antena y, por tanto, en mejores condiciones de cobertura que el usuario emulado UE0.

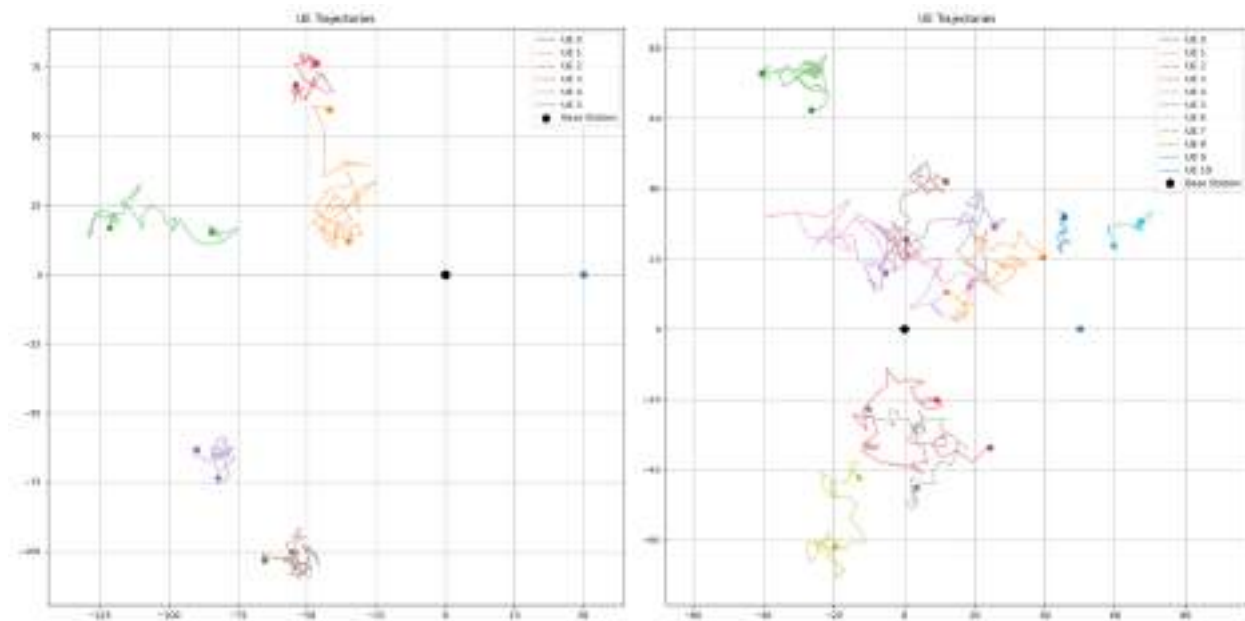


Fig. 3.13. Trayectorias en banda n40 (izquierda) y n78 (derecha)

La siguiente figura muestra los histogramas de SINR de dichas pruebas. Se puede comprobar la dispersión tan alta de los UEs en movimiento, que muestra la alta variabilidad del canal.

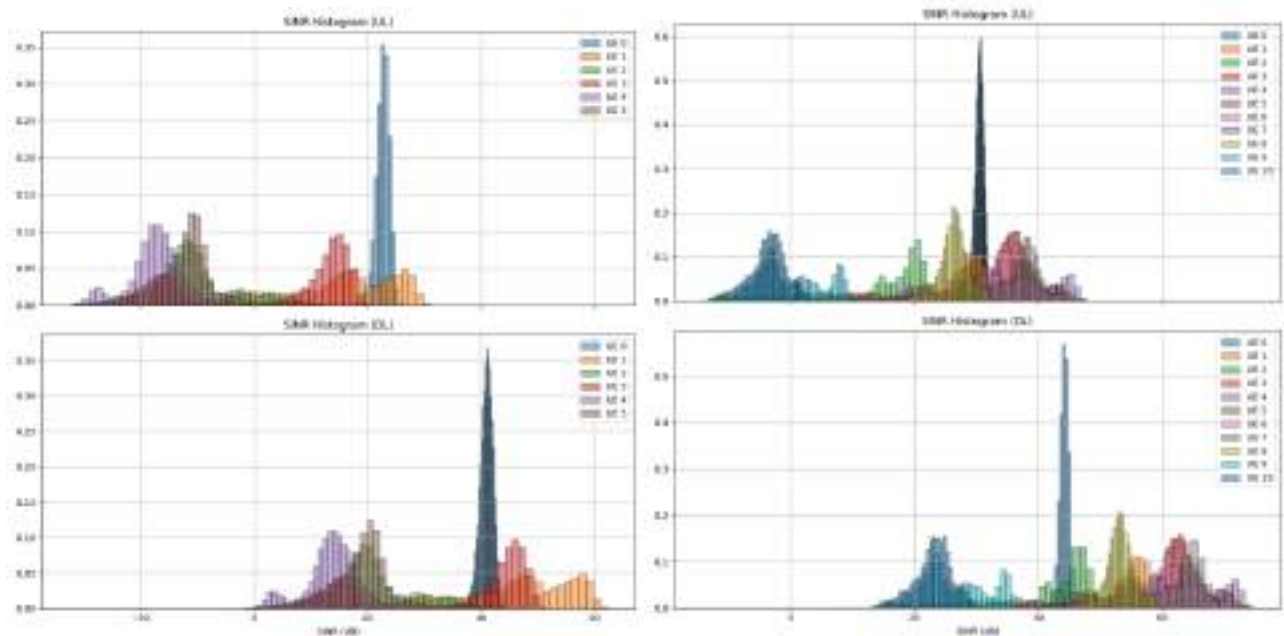
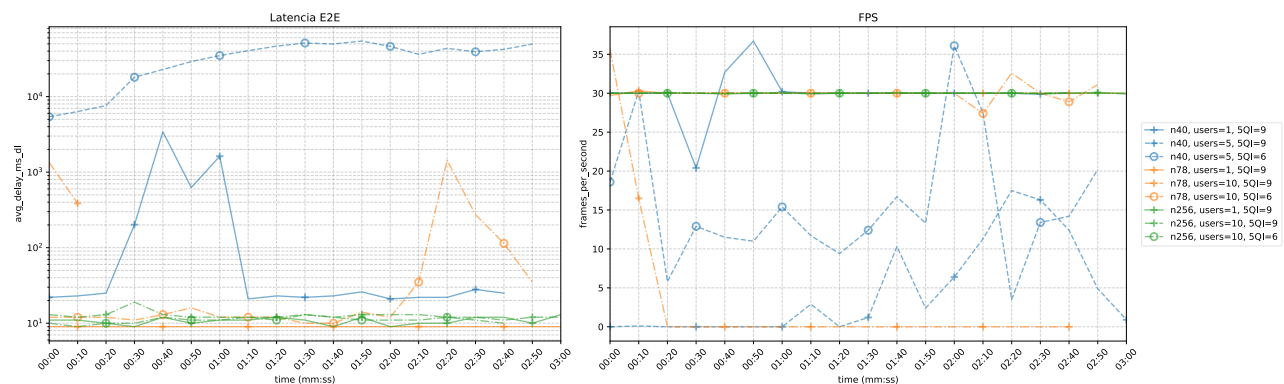


Fig. 3.14. Histogramas de SINR en banda n40 (izquierda) y n78 (derecha)

La siguiente figura muestra las gráficas de distintos KPIs de aplicación. Se ha usado una resolución de referencia de 1280x480 píxeles para las imágenes a 30 FPS, lo que se traduce en unos 10.5 Mbps. Se observa que la celda en n40 tiene problemas incluso para gestionar un único usuario, o para dar prioridad usando únicamente 5QIs; que la celda en n78 presta servicio con suficiente QoS con 10 usuarios concurrentes siempre y cuando se priorice al usuario emulado; y que la celda en banda milimétrica n258, de mucho mayor ancho de banda, tiene capacidad para todo el tráfico incluso sin priorizar. En situaciones de carga en las que se presta servicio, como puede ser n40 con un usuario adicional o n78 con 5QI=6, se observa que el servicio se presta, pero no de forma estable (hay picos en los FPSs y en la latencia).



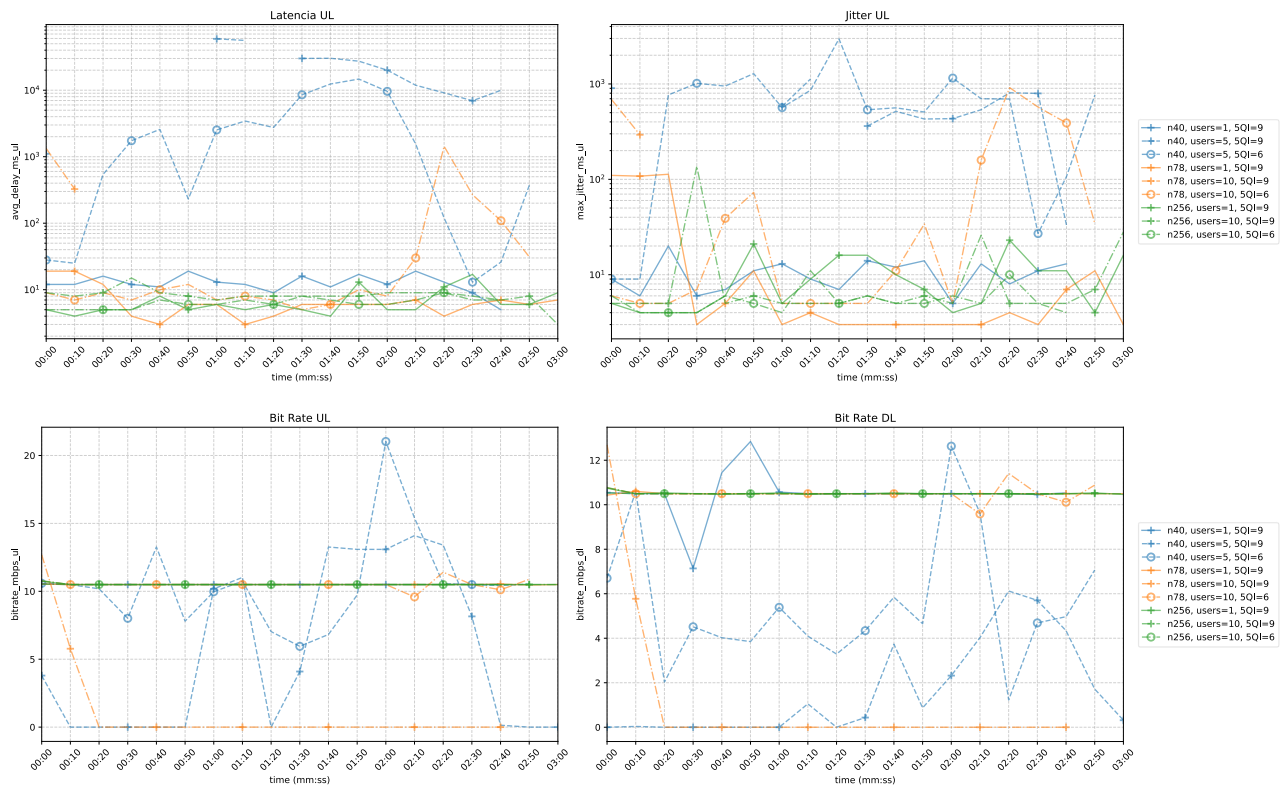


Fig. 3.15. Gráficas de KPIs de aplicación para los distintos escenarios de computación delegada

A partir de estas métricas se estima el MOS del servicio, teniendo en cuenta tanto la latencia total, a la que se añade un tiempo estimado de inferencia en el MEC de 13ms, como la resolución del vídeo. Este análisis permite identificar el impacto del número de usuarios y de la banda de frecuencia en la viabilidad del escenario.

Tabla 9. MOS estimado para el escenario de procesamiento delegado en función de la banda y el número de usuarios

	Escenario	n40			n78			n256		
Usuarios concurrentes	1-10	1	5	5	1	10	10	1	10	10
5QI		9	9	6	9	9	6	9	9	6
MOS	4.0	3.8	1.0	1.6	4.6	1.6	4.1	4.6	4.6	4.6
DL Rate	1-20 Mbps	10.5	2.0	5.5	10.5	1.1	10.5	10.5	10.5	10.5
UL Rate	10-200 Mbps	10.5	4.3	10.5	10.5	1.1	10.5	10.5	10.5	10.5
FPS	30	30.0	5.7	15.7	30.0	3.1	30.0	30.0	30.0	30.0

UL PDB	PDB < 20 ms	12.8	28000	3400	7.0	830	110	6.8	8.3	6.4
E2E	RTT < 40 ms	350	-	34000	9.0	860	114	10.8	12.7	11.0

3.3.3 Escenario 3: Regulación Emocional (*Emotional Regulation*)

El escenario de regulación emocional se caracteriza por tasas binarias reducidas y una menor sensibilidad a la latencia en comparación con los escenarios anteriores. Las pruebas se realizan en un entorno de macrocelda urbana en banda n40, con un número elevado de usuarios concurrentes.

Tabla 10. Parámetros de simulación del escenario de regulación emocional.

Macrocelda Urbana	Escenario E1.2	Simulación
Banda		n40
Ancho de Banda		20 MHz TDD
Distancia a la antena	500 m	200m
DL Rate	0.1 Mbps	0.13
UL Rate	0.1 Mbps	0.13
Tipo Tráfico	A ráfagas	10 muestras/s
Usuarios concurrentes	10-50	1, 50
Distancia media de los usuarios	1000 m	500m
Distancia máxima de los usuarios	5000 m	2500m
DL Rate	0.1 Mbps	0.1
UL Rate	0.1 Mbps	0.1
Tipo Tráfico	A ráfagas	

La principal característica de este escenario es que se puede simular un número alto de usuarios concurrentes (50). La siguiente figura muestra su distribución espacial y la variedad de los histogramas de SINR.

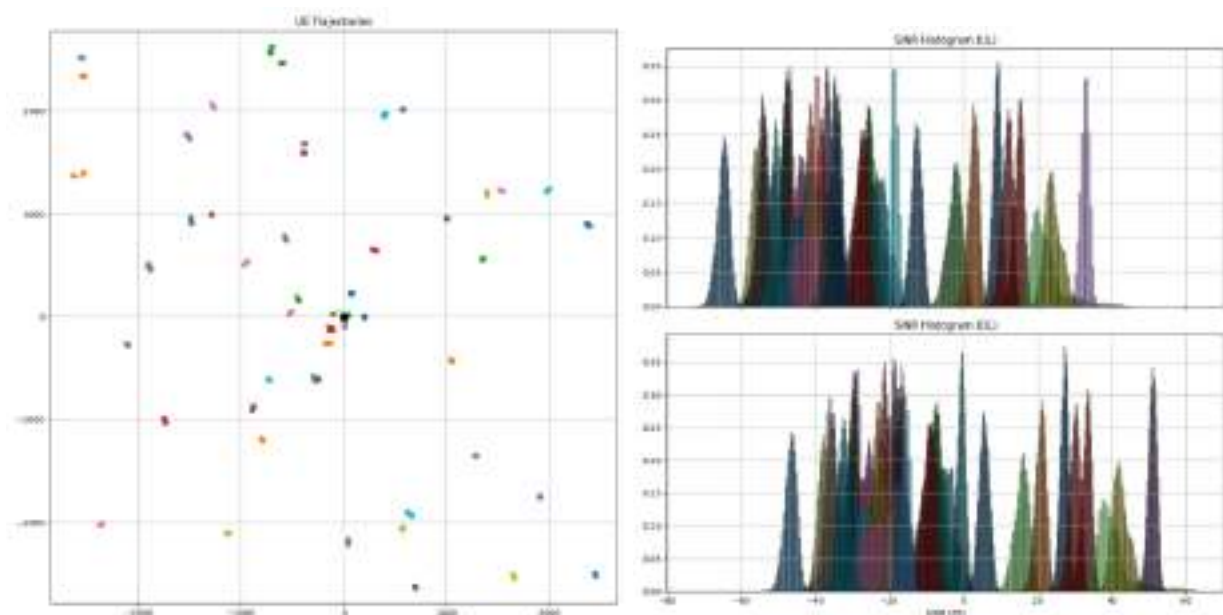


Fig. 3.16. Distribución de UEs e histogramas de SINR para el escenario de regulación emocional

Las gráficas siguientes muestran la evolución de la tasa de muestreo, latencia y bitrate durante la prueba, para las tres condiciones (1 y 50 usuarios concurrentes, con y sin priorización del usuario emulado).

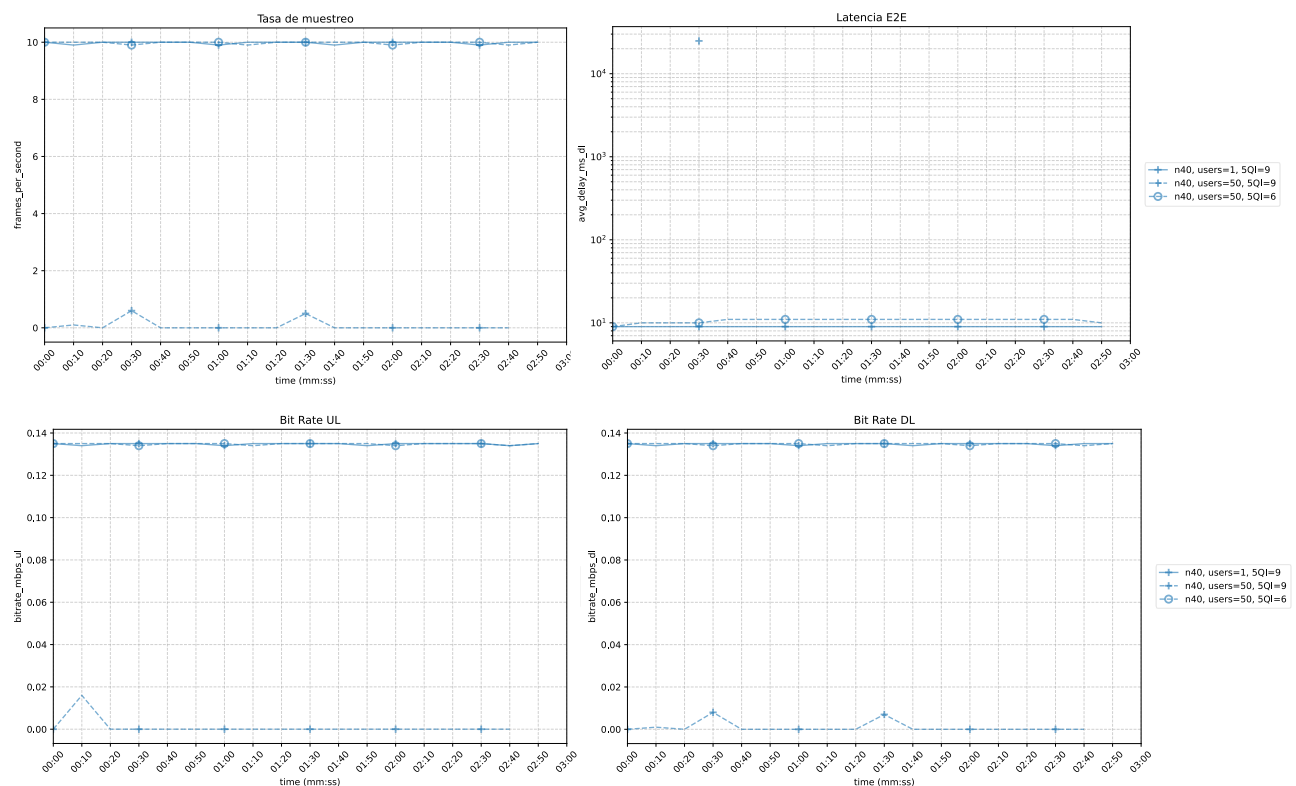


Fig. 3.17. KPIs para el escenario de regulación emocional

La siguiente tabla muestra un resumen de los KPIs, así como el MOS calculado con el modelo paramétrico, asumiendo un tiempo de inferencia de 10 ms. Se comprueba que, para escenarios con muchos usuarios concurrentes, es preciso utilizar gestión de QoS para poder proporcionar suficiente QoE en este escenario de regulación emocional.

Tabla 11. MOS estimado para el escenario de regulación emocional

	Escenario	n40		
Usuarios concurrentes	1-10	1	5	5
5QI		9	9	6
MOS		4.7	1.0	4.7
DL Rate	100 kbps	135	0.1	135
UL Rate	100 kbps	135	0.1	135
FPS	10	10.0	0.07	10.0
UL delay	<300 ms PDB	4.0	-	6.7
E2E delay	<600 ms RTT	9.0	-	10.1

3.4

Pruebas en red física

En este apartado se presentan los resultados de las pruebas realizadas sobre la red 5G física desplegada en la Fundación Juan XXIII, utilizando mecanismos de *network slicing* para controlar la calidad de servicio. A diferencia del apartado anterior, basado en red emulada, estas pruebas se han llevado a cabo en un entorno real, con equipos comerciales y tráfico generado sobre la infraestructura 5G en banda n40 descrita en el capítulo 2.

Como se describe en el entregable E2.2, la principal herramienta utilizada en el proyecto para garantizar calidad de servicio (QoS) en la red 5G para dar soporte a los distintos casos de uso del proyecto es el network slicing. Para el proyecto, se han definido tres slices: una de PRIORITY (Gold) enfocada al tráfico principal de uplink, pero con suficiente capacidad de downlink; una para HMDs (Silver), que utiliza fundamentalmente el downlink de la red; y una DEFAULT (Bronze) para el tráfico que no pertenece a los casos de uso (mantenimiento, etc.).

El objetivo de estas pruebas es evaluar hasta qué punto el *slicing* permite proteger el tráfico asociado a los escenarios de aplicación del proyecto frente a tráfico competitivo, y analizar el impacto de dicha protección tanto en métricas de red como en la calidad de experiencia estimada.

Tabla 12. Configuración de slices de red

Parameter	PRIORITY (Golden)	HMDS (Silver)	DEFAULT (Bronze)
UL minQuota	50	15	5
UL maxQuota	100	20	20
DL minQuota	15	50	5
DL maxQuota	20	100	20

Para la campaña de pruebas, se han usado dos UEs: uno para representar el escenario de aplicación bajo estudio y otro para modelar el tráfico que compite con él. Para todos los escenarios se emplea una configuración similar: un equipo de usuario (UE A) que representa el caso de uso bajo estudio y un segundo equipo (UE B) que genera tráfico competitivo. Se comparan resultados con y sin *slicing*, o entre distintas configuraciones de *slice*, manteniendo constantes el resto de parámetros. La siguiente tabla resume el contexto de todas las pruebas.

Tabla 13. Configuración de la prueba

UE		Localización	Tipo de tráfico	Slice (Prueba 1)	Slice (prueba 2)
A	1	Atención Psicológica	Escenario de aplicación	Sin slicing	PRIORITY
B	5	Aula de Informática	Ruido (competitivo)	Sin slicing	HMDS

3.4.1 Escenario 1: Teleportación Virtual (Virtual Teleportation)

Este escenario evalúa el comportamiento del servicio de teleportación virtual en un entorno real cuando existe tráfico competitivo en el enlace ascendente, que es el más crítico para este caso de uso. El tráfico generado por el UE A simula la transmisión de vídeo inmersivo a distintos bitrates, mientras que el UE B introduce carga adicional mediante tráfico sintético.

Para las pruebas en el Escenario 1, se empleó la configuración de la siguiente tabla.

Tabla 14. Configuración de la prueba del Escenario 1

UE	Localización	Tipo de tráfico	Slice	Slice
----	--------------	-----------------	-------	-------

				(Prueba 1)	(prueba 2)
A	1	Atención Psicológica	Tráfico de vídeo (Snowl) UL: 5, 7.5, 10, 12.5, 15, 17.5, 20 Mbps	Sin slicing	PRIORITY
B	5	Aula de Informática	Tráfico iperf3 UL (cada minuto): 0, 5, 10, 15, 20 Mbps (prueba 1) 0, 3, 6 Mbps (prueba 2)	Sin slicing	HMDS

Se simuló tráfico de uplink procedente de un búho (vídeo inmersivo) a distintos bitrates desde el UE A. Para cada uno de ellos, se generó tráfico UDP de uplink con iperf3 desde el UE B, partiendo de cero e incrementando en 5 Mbps cada minuto. Se midió el tráfico recibido en el core de la red, así como la recepción del vídeo en RTP en el backend del búho instalado en el MEC. Para las pruebas con slicing, dado que la capacidad máxima de uplink para el UE B está en torno a 4 Mbps (un 20% de 20 Mbps), el tráfico de ruido generado es menor.

3.4.1.1 Sin configuración de slicing

En primer lugar, se analizan los resultados obtenidos en el escenario de teleportación virtual cuando ambos UEs comparten los recursos de la red sin ningún mecanismo de priorización. El UE A transmite tráfico de vídeo inmersivo en enlace ascendente a distintos bitrates, mientras que el UE B genera tráfico competitivo creciente.

La Figura 3.18 muestra la evolución temporal del tráfico de uplink recibido en el core de red para ambos UEs. Se observa que, a medida que aumenta el tráfico competitivo, el uplink total de la celda se aproxima rápidamente a su capacidad máxima, produciéndose una saturación del enlace.



Fig. 3.18. Tráfico de uplink recibido en el core desde los UEs A (rojo) y B (verde)

Esta saturación tiene un impacto directo sobre el rendimiento del vídeo transmitido por el UE A. Como se aprecia en la Figura 3.19, el bitrate efectivo recibido en el backend deja de crecer a partir de ciertos valores, independientemente del bitrate configurado en el origen.



Fig. 3.19. Medidas en el backend: bit rate de video (recibido), tasa de pérdida de paquetes y retardo medio en de paquete (one-way)

La Tabla 14 resume los valores medios medidos para cada configuración de bitrate. Para bitrates de vídeo de hasta 10 Mbps, el sistema mantiene pérdidas despreciables y latencias reducidas. Sin embargo, a partir de 12.5 Mbps aparecen incrementos significativos de latencia, y para bitrates de 15 Mbps o superiores se observan tasas de pérdida elevadas y retardos del orden de varios cientos de milisegundos.

Tabla 15. Resumen del Escenario 1, sin slicing, con tráfico competitivo de 20 Mbps.

UE A			UE B	
UE	Core	Backend	UE	Core

Video Bitrate (Mbps)	Uplink Bitrate (Mbps)	Video Bitrate (Mbps)	Packet Loss Rate (%)	Packet Delay (Ms)	Noise Bitrate (Mbps)	Uplink Bitrate (Mbps)
5.0	5.3	5.2	0	15	20	14.2
7.5	7.9	7.8	0	15	20	10.1
10.0	10.5	10.3	0	15	20	11.2
12.5	13.2	12.8	0	35	20	9.2
15.0	14.3	13.9	9.3	310	20	11.6
17.5	14.3	13.9	22	480	20	11.7
20.0	14.3	14.0	32	760	20	11.6

Estos resultados indican que, en ausencia de mecanismos de priorización, el tráfico competitivo en uplink limita severamente la capacidad disponible para el servicio de teleportación virtual, impidiendo escalar el bitrate del vídeo y degradando notablemente la calidad percibida.

3.4.1.2 Con configuración de slicing

A continuación, se repite el experimento activando network slicing, asignando el tráfico de vídeo del UE A al slice PRIORITY y el tráfico competitivo del UE B al slice HMDS, con los límites de capacidad definidos en la configuración de red.

La siguiente figura muestra la monitorización del tráfico de uplink recibido en el core bajo esta configuración. En este caso, el tráfico competitivo queda claramente acotado, mientras que el tráfico asociado al vídeo dispone de recursos suficientes incluso para los bitrates más elevados.



Fig. 3.20. Tráfico de uplink recibido en el core desde los UEs A (rojo) y B (verde)

El impacto de esta separación de recursos se refleja directamente en las métricas de aplicación. Como se observa en la siguiente figura, el bitrate recibido en el backend sigue de forma fiel al bitrate transmitido por el UE A, sin apreciarse efectos de saturación ni degradación progresiva.



Fig. 3.21. Medidas en el backend: bit rate de video (recibido), tasa de pérdida de paquetes y retardo medio en de paquete (one-way)

La siguiente tabla recoge los valores medios medidos para cada bitrate de vídeo. En todas las configuraciones evaluadas, incluyendo bitrates de hasta 20 Mbps, se mantienen pérdidas nulas y latencias estables en torno a unas pocas decenas de milisegundos.

Tabla 16. Resumen del Escenario 1, con slicing, con tráfico competitivo de 6Mbps.

UE A			UE B	
UE	Core	Backend	UE	Core

Video Bitrate (Mbps)	Uplink Bitrate (Mbps)	Video Bitrate (Mbps)	Packet Loss Rate (%)	Packet Delay (Ms)	Noise Bitrate (Mbps)	Uplink Bitrate (Mbps)
5	5.3	5.2	0	12	6	3.9
7.5	7.9	7.8	0	12	6	3.9
10	10.5	10.3	0	12	6	3.7
12.5	13.2	12.8	0	12	6	3.9
15	15.8	15.4	0	12	6	3.9
17.5	18.4	18.0	0	12	6	3.5
20	21.0	20.5	0	12	6	3.0

Estos resultados confirman que el uso de *network slicing* permite aislar eficazmente el tráfico crítico del escenario de teleportación virtual frente al tráfico competitivo, garantizando un comportamiento estable y predecible del servicio.

3.4.1.3 Comparación

Para facilitar la comparación directa entre ambos escenarios, la siguiente figura muestra el bitrate efectivo recibido en backend en función del bitrate transmitido, con y sin slicing, para la prueba en la que se envía video a 20 Mbps. Sin priorización, el bitrate útil se estanca alrededor de 14 Mbps, mientras que con slicing escala de forma prácticamente lineal hasta los valores máximos configurados.



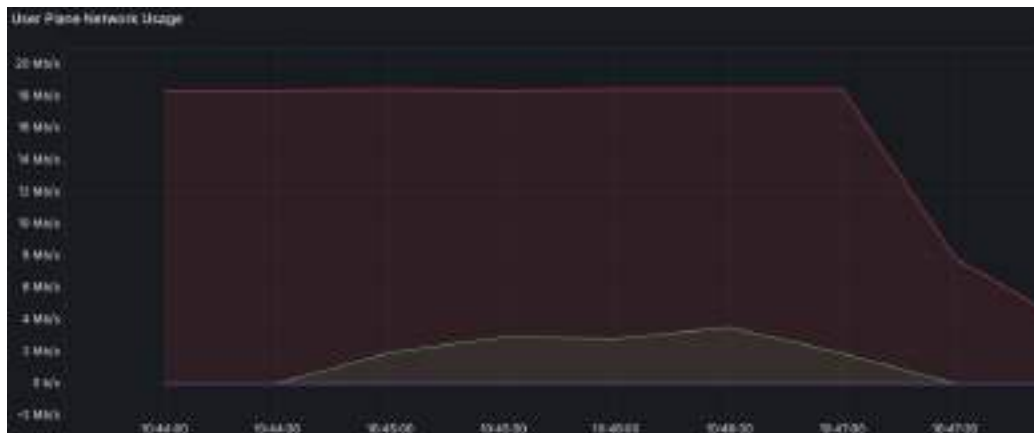


Fig. 3.22. Detalle del uplink de los UES A (rojo, escenario de prueba) y B (verde, ruido) para el caso de prueba de 20 Mbps de bit rate de video. Sin slicing en la parte superior, con slicing en la inferior.

La siguiente figura muestra el bitrate medido en capa de aplicación, que se corresponde con el tráfico del UE A de la imagen anterior.



Fig. 3.23. Comparación del bit rate recibido sin (izquierda) o con (derecha) el slicing activo

La siguiente figura presenta la tasa de pérdidas de paquetes medida en ambos casos. En ausencia de slicing, las pérdidas aumentan rápidamente a partir de 15 Mbps, alcanzando valores elevados, mientras que con slicing se mantienen en cero para todo el rango de bitrates evaluado.

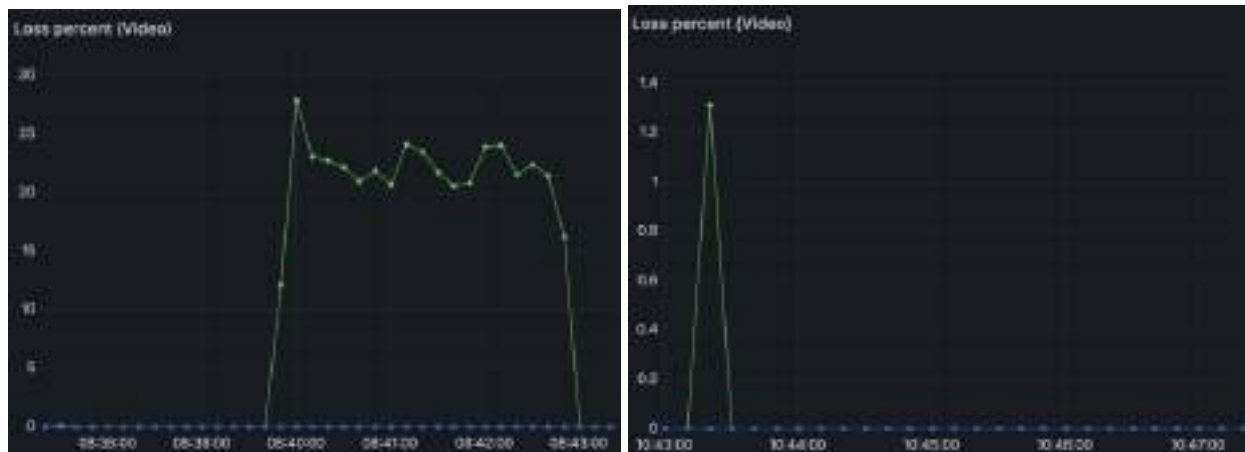


Fig. 3.24. Comparación de la tasa de pérdidas sin (izquierda) o con (derecha) el slicing activo

De forma análoga, los siguientes gráficos muestran la latencia media extremo a extremo. Sin slicing se observan incrementos muy significativos bajo carga, con retardos del orden de cientos de milisegundos, mientras que con slicing la latencia permanece baja y prácticamente constante.

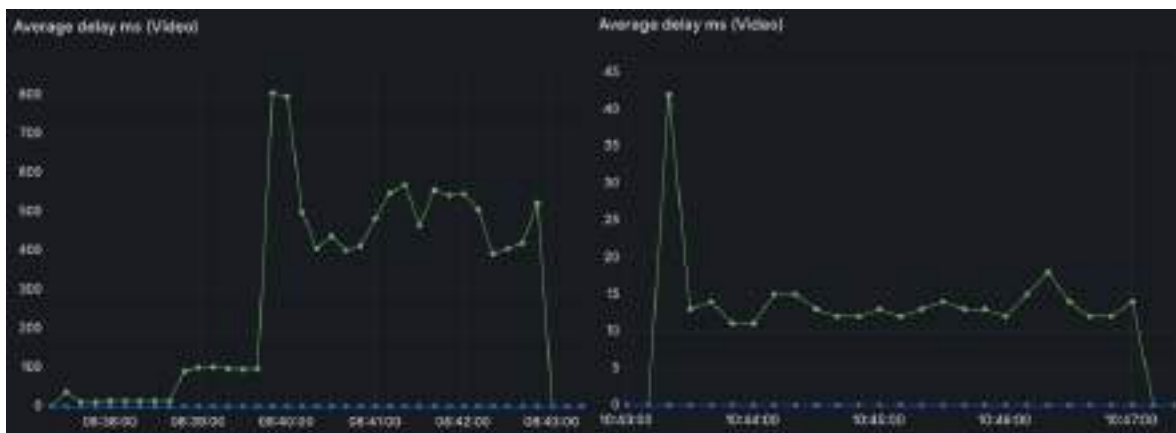


Fig. 3.25. Comparación de la latencia sin (izquierda) o con (derecha) el slicing activo

3.4.1.4 Análisis de QoE

Finalmente, a partir de las métricas medidas se estima la calidad de experiencia del servicio mediante el modelo paramétrico definido en el apartado 3.1. La figura muestra la evolución del MOS estimado en función del bitrate de vídeo. Sin slicing, el MOS cae por debajo de

valores aceptables para bitrates altos, mientras que con slicing se mantiene por encima de MOS = 4 en todo el rango considerado.

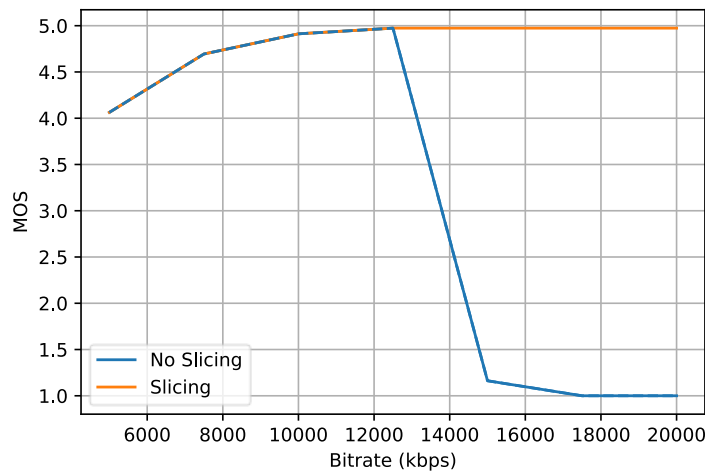


Fig. 3.26. MOS estimado en función del bitrate transmitido desde A, asumiendo que desde B se genera simultáneamente 20 Mbps

En conjunto, estos resultados evidencian que la protección del uplink mediante *network slicing* no solo mejora las métricas de red, sino que resulta imprescindible para garantizar la viabilidad del servicio de teleportación virtual en un entorno real con tráfico competitivo.

3.4.2 Escenario 2: Procesamiento delegado de Realidad Mixta (*Mixed Reality computation offloading*)

En este escenario se evalúa el funcionamiento del procesamiento delegado en un entorno real, considerando tráfico simétrico de subida y bajada y una elevada sensibilidad a la latencia. El UE A genera tráfico asociado a la transmisión de imágenes para segmentación en el MEC, mientras que el UE B introduce tráfico competitivo en el enlace descendente.

Las pruebas se han realizado para distintas resoluciones de imagen, manteniendo una tasa de cuadros constante, con el objetivo de analizar el impacto de la resolución sobre el bitrate, la latencia y la estabilidad del servicio en presencia de tráfico competitivo.

Para las pruebas en el Escenario 2, se empleó la configuración de la siguiente tabla.

Tabla 17. Configuración de la prueba del Escenario 2

UE	Localización		Tipo de tráfico	Slice
A	1	Atención Psicológica	Tráfico de segmentación, simétrico Resoluciones: 1080, 810, 720, 540	PRIORITY

B	5	Aula de Informática	Tráfico iperf3 DL (cada minuto): 10, 50, 100 Mbps	HMDS
----------	---	---------------------	--	------

Se simuló tráfico de uplink para computación delegada (imágenes JPEG) a distintas resoluciones desde el UE A. Para gestionar un caso peor, se reenvía el mismo frame en el downlink (tráfico simétrico). Para cada una de las resoluciones, se generó tráfico TCP de downlink con iperf3 hacia el UE B. Se midió el tráfico recibido en el core de la red, así como en el backend de *offloading* instalado en el MEC. En este caso, se prueba únicamente con slicing, ya que la capacidad asimétrica de la red hace que el cuello de botella sea el uplink total que la red es capaz de gestionar.

La siguiente figura muestra el tráfico de enlace descendente generado en el core durante las pruebas. Se observa que el volumen de tráfico aumenta de forma proporcional con la resolución de las imágenes transmitidas, situándose en los rangos esperados para cada configuración.

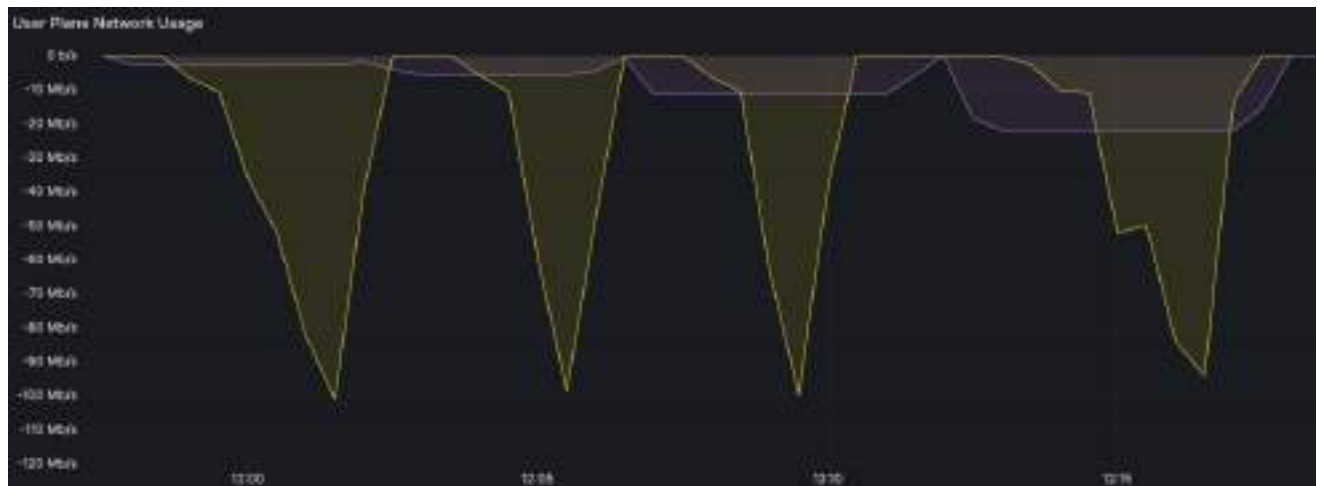


Fig. 3.27. Tráfico de downlink enviado desde el core a los UEs A (morado) y B (amarillo)

Este aumento de tráfico tiene un efecto directo sobre las métricas observadas en el backend de la aplicación. En la figura siguiente se presentan el retardo extremo a extremo y la tasa de cuadros efectiva para las distintas resoluciones evaluadas. Para resoluciones bajas e intermedias, el sistema mantiene valores estables de latencia y una tasa de cuadros cercana al objetivo, mientras que resoluciones más exigentes introducen mayores retardos y una mayor variabilidad en el servicio.



Fig. 3.28. Medidas en el backend: bit rate de video (recibido), tasa de paquetes por segundo (Frame Rate) y retardo medio en de paquete (one-way)

La tabla siguiente resume los valores medios obtenidos para cada configuración de resolución. En ella se aprecia que, a medida que aumenta la resolución, el retardo total se incrementa debido tanto al mayor volumen de datos transmitidos como al aumento del tiempo de procesamiento en el MEC. Este efecto limita la estabilidad del servicio para las configuraciones más exigentes.

Tabla 18. Resumen del Escenario 2, con slicing, con tráfico competitivo de 100 Mbps de downlink

UE A							UE B	
UE			Core	Backend			Core	UE
Resolución	FPS	Bitrate vídeo	DL	BR	FPS	Delay	Bitrate ruido	DL noise
960x540	30	3	2.9	2.7	30	11	100	101

1280x720	30	5	5.7	5.3	30	14	100	99
1440x810	30	10	11.1	10.5	30	17	100	100
1920x1080	30	21	21.9	20.7	30	35	100	94

3.4.2.1 Análisis de QoE

A partir de las métricas medidas, se estima la calidad de experiencia del servicio utilizando el modelo paramétrico definido en el apartado 3.1. La siguiente figura muestra el valor de MOS estimado en función de la resolución de imagen. Los resultados indican que las resoluciones de trabajo consideradas en el proyecto permiten alcanzar valores de MOS compatibles con una experiencia de usuario satisfactoria, mientras que resoluciones más altas penalizan la QoE debido al incremento del retardo extremo a extremo.

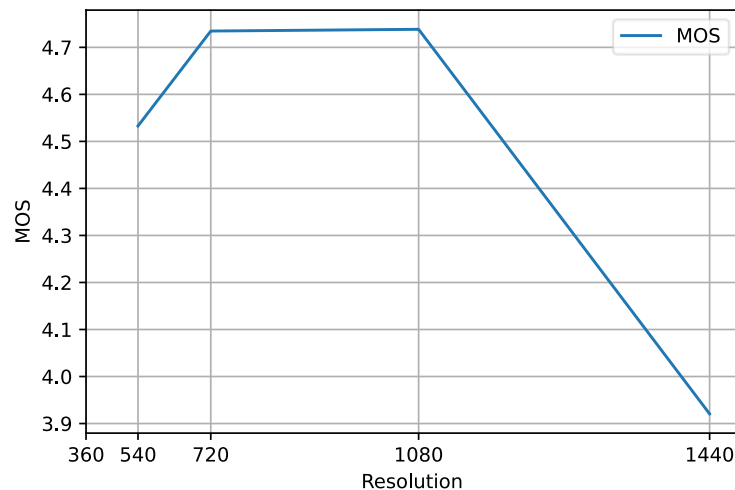


Fig. 3.29. MOS estimado en función de la resolución

En conjunto, los resultados confirman que el escenario de procesamiento delegado de realidad mixta puede ejecutarse de forma estable sobre la red 5G física desplegada, siempre que se mantenga un compromiso adecuado entre resolución, carga de red y capacidad de procesamiento en el edge, y que los mecanismos de control de calidad permiten sostener una experiencia consistente en condiciones reales.

3.4.3 Escenario 3: Regulación Emocional (*Emotional Regulation*)

En este escenario se evalúa el comportamiento del servicio de regulación emocional en un entorno real, caracterizado por tasas binarias reducidas y una menor exigencia de ancho de

banda en comparación con los escenarios anteriores. No obstante, el servicio presenta requisitos de latencia suficientemente estrictos para garantizar una respuesta adecuada del sistema en tiempo casi real.

El UE A genera tráfico asociado a la transmisión periódica de datos y resultados de inferencia, mientras que el UE B introduce tráfico competitivo de mayor volumen. Las pruebas se han realizado con *network slicing* activado, asignando el tráfico del escenario de aplicación a un slice prioritario.

La siguiente figura muestra el tráfico de enlace descendente observado en el core durante la ejecución de las pruebas. Se aprecia que el tráfico asociado al servicio de regulación emocional mantiene un comportamiento estable, incluso en presencia de tráfico competitivo significativo generado por el segundo UE.



Fig. 3.30. Tráfico de downlink enviado desde el core a los UEs A (morado) y B (amarillo)

Este comportamiento se refleja en las métricas observadas en el backend de la aplicación. En la figura siguiente se presenta la evolución temporal del bitrate, la latencia extremo a extremo y la tasa de muestreo efectiva. Los resultados muestran que el servicio mantiene valores estables de latencia y una tasa de muestreo acorde a los requisitos del escenario.

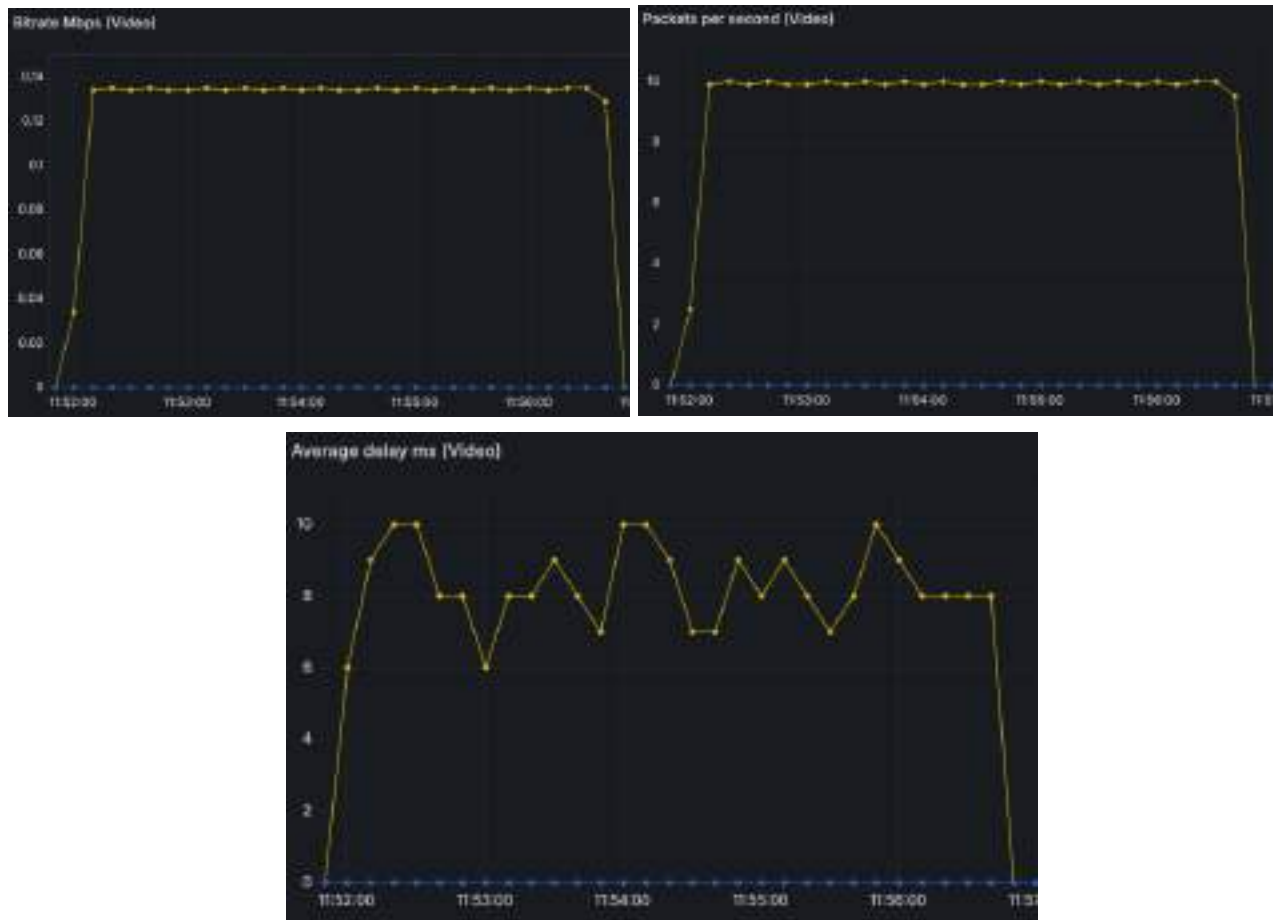


Fig. 3.31. Medidas en el backend: bit rate de datos (recibido), tasa de paquetes por segundo (Sample Rate) y retardo medio en de paquete (one-way)

La tabla siguiente resume los valores medios medidos durante la prueba. Los resultados confirman que, gracias a la priorización del tráfico, el servicio cumple holgadamente los requisitos de retardo definidos para este escenario, sin verse afectado de forma significativa por la carga introducida por el tráfico competitivo.

Tabla 19. Resumen del Escenario 3, con slicing, con tráfico competitivo de 100 Mbps de bandwidth

UE A							UE B	
UE		Core	Backend				Core	UE
Sample Rate	Bitrate datos	DL	BR	Sample Rate	Delay	MOS	Bitrate ruido	DL noise
10	0.15	0.15	0.15	10	8	4.6	100	100

A partir de las métricas obtenidas se estima la calidad de experiencia del servicio utilizando el modelo paramétrico definido en el apartado 3.1. Los valores de MOS estimados se sitúan en un rango elevado y estable, coherente con las expectativas para un servicio con requisitos moderados de ancho de banda y latencia.

En conjunto, los resultados indican que el escenario de regulación emocional puede desplegarse de forma robusta sobre la red 5G física, y que el uso de *network slicing* permite garantizar la estabilidad del servicio incluso en condiciones de carga elevada, cumpliendo los objetivos de calidad definidos para este caso de uso.

Conclusiones

En este capítulo se han validado los modelos y técnicas empleados para modelar y controlar la calidad de servicio y la calidad de experiencia en los tres escenarios del proyecto. Los modelos paramétricos de QoE definidos permiten traducir métricas objetivas de red en estimaciones coherentes de experiencia de usuario, y capturan de forma adecuada los factores dominantes de cada caso de uso: el bitrate y la tasa de pérdidas en el escenario de teleportación virtual, la latencia y la resolución en el procesamiento delegado de realidad mixta, y la latencia extremo a extremo en el escenario de regulación emocional.

Las pruebas realizadas sobre red emulada muestran que, en ausencia de mecanismos de gestión de QoS, el aumento de la carga de usuarios concurrentes degrada rápidamente la calidad de los servicios más exigentes. En este contexto, la asignación de prioridades mediante 5QIs permite proteger de forma efectiva el tráfico crítico, manteniendo valores de MOS acordes a los objetivos definidos en el proyecto. De forma complementaria, los resultados obtenidos en red física confirman que el comportamiento observado en red emulada es representativo de un despliegue real.

Las pruebas sobre la red 5G física evidencian que el uso de *network slicing* no solo mejora las métricas de rendimiento de la red, sino que resulta clave para habilitar de forma efectiva los casos de uso del proyecto en presencia de tráfico competitivo. La combinación de modelos de QoE, experimentación en red emulada y validación en red física proporciona así una metodología consistente para la evaluación de servicios XR sobre redes 5G.

En conjunto, la validación realizada confirma que los requisitos definidos en el entregable E1.2 son alcanzables bajo configuraciones de red adecuadas y con mecanismos de priorización activos..



4 Pruebas finales del sistema de edge y de offloading de IA

Este capítulo presenta la validación del sistema de edge y de offloading de IA desplegado como parte de la infraestructura del proyecto. En primer lugar, se describe la arquitectura implementada y se comprueba el correcto despliegue, configuración y estado operativo de los distintos componentes del sistema. A continuación, se evalúa el rendimiento del nodo MEC mediante pruebas de carga y ejecución bajo distintas condiciones. Los resultados obtenidos permiten analizar la viabilidad del sistema de edge como soporte de los escenarios de aplicación del proyecto.

Pruebas funcionales

4.1

4.1.1 Arquitectura general

Nokia OpenEdge es una solución integral de computación en el borde diseñada para acercar las capacidades de la nube a los usuarios finales, reduciendo así la latencia y mejorando el rendimiento del servicio. Soporta una amplia gama de aplicaciones, incluyendo IoT, automatización industrial y experiencias inmersivas.

Para el proyecto se ha desplegado un chasis OpenEdge con alimentación en continua y dos servidores físicos:

- **Servidor Cabezon:** Un servidor open-edge de 1U diseñado para tareas generales de computación en el borde. Es compacto y adecuado para su despliegue en entornos con limitaciones de espacio.
- **Servidor Casals:** Un servidor open-edge de 2U equipado con 2 GPUs Tesla T4, proporcionando capacidades computacionales mejoradas para aplicaciones de IA y aprendizaje automático.

Cada uno de los servidores físicos tiene su propio entorno de virtualización PROXMOX 8.0.3 para gestionar máquinas virtuales (VMs) y dispositivos de almacenamiento. Las VM están configuradas para manejar varias funciones de red y aplicaciones, mientras que los dispositivos de almacenamiento proporcionan las capacidades necesarias de almacenamiento de datos y respaldo.

El servidor Cabezon contiene y administra los siguientes hosts virtuales:

- 100 (camarlengo). VM responsable de la administración del MEC. Contiene los servicios principales del nodo: servidor DNS, servidor openVPN, pasarela SSH.

- 101 (Urbano2) / 500 (madf-core-sa1) / 501 (madf-core-sa1). VM principal del core de red 5G (cada una contiene una versión de configuración diferente)
- 102 (Benedicto16). VNF de monitorización de KPIs. Contiene el stack TIG (telegraf, influxDB, grafana).
- 103 (Gregorio15). VNF de gestión de contenidos (CMS). Contiene un servidor MediaCMS modificado para almacenar los contenidos inmersivos de los distintos casos de uso.

El servidor Casals contiene y administra los siguientes hosts virtuales:

- 101 (Clemente10). VNF de offloading de IA. Accede a una de las GPUs NVidia Tesla T4 disponibles en el OpenEdge, para la aceleración de la ejecución de algoritmos IA.
- 102 (Clemente11). VNF de offloading de IA. Accede la otra GPUs NVidia Tesla T4 disponible en el OpenEdge, para la aceleración de la ejecución de algoritmos IA.
- 103 (Leon13). VNF de control de las aplicaciones de terapia. Esta VNF ejecuta el backend de las aplicaciones de terapia o formación. Se trata de un servidor basado en Node.js, que usa un servidor Express para servir las páginas de operación y configuración, y Socket.IO para la mensajería.
- 104 (Pio12). VNF de backend de telepresencia. Esta VNF contiene el backend del sistema de telepresencia (Snowl), tanto en plano de datos como en plano de control.

Dado que las VNFs realizan funciones complejas, el software de servicio está empaquetado en contenedores docker. De este modo, es posible reconfigurar o arrancar internamente distintas versiones del VNF arrancando y parando distintos contenedores. Esto se aplica a las VNFs de servicio: Clemente10, Clemente11, Pio12, Leon13 y Gregorio15.

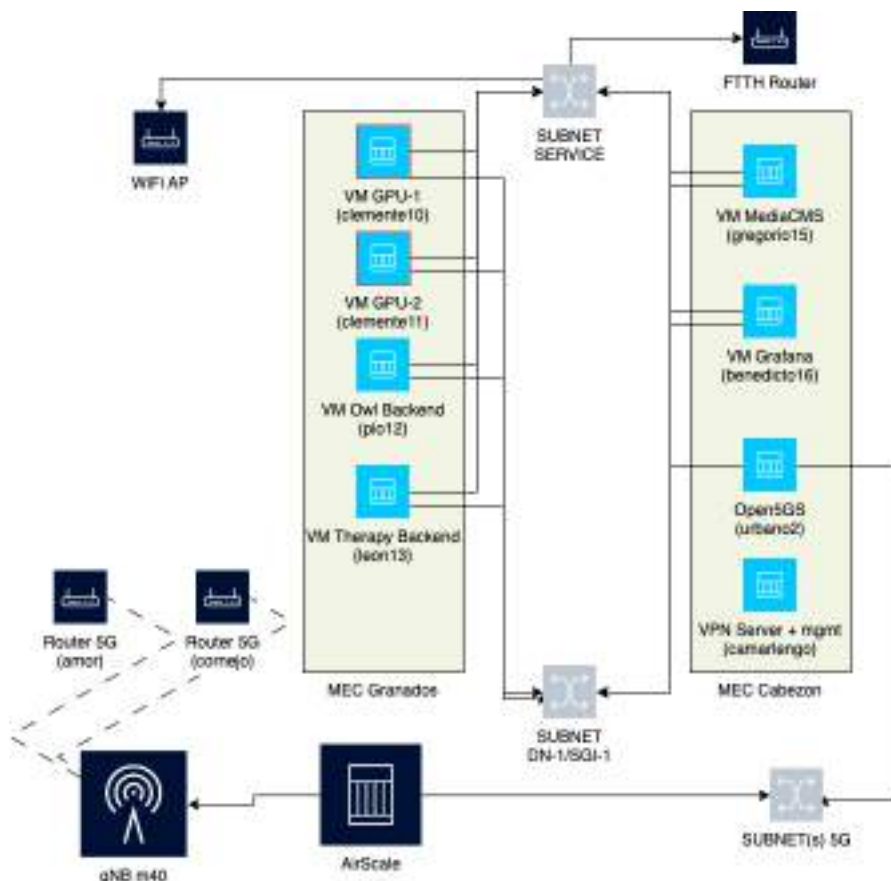


Fig. 4.1. Arquitectura general del Sistema Edge

4.1.2 Sistema de virtualización

El sistema de edge desplegado se apoya en un entorno de virtualización basado en Proxmox, que permite la gestión centralizada de los recursos de computación, almacenamiento y red de los servidores físicos que componen el nodo MEC. Este entorno constituye la base sobre la que se ejecutan las distintas máquinas virtuales que implementan las funciones de red y los servicios de aplicación del proyecto.

En particular, cada uno de los servidores físicos OpenEdge dispone de su propio entorno Proxmox, lo que facilita la separación de roles y la asignación diferenciada de recursos. En el servidor Cabezon, orientado a tareas generales de computación y servicios auxiliares del MEC, se alojan las máquinas virtuales relacionadas con la administración del nodo, el core 5G y la monitorización. La consola de gestión correspondiente permite visualizar el estado operativo de las VMs, su consumo de recursos y la conectividad asociada.

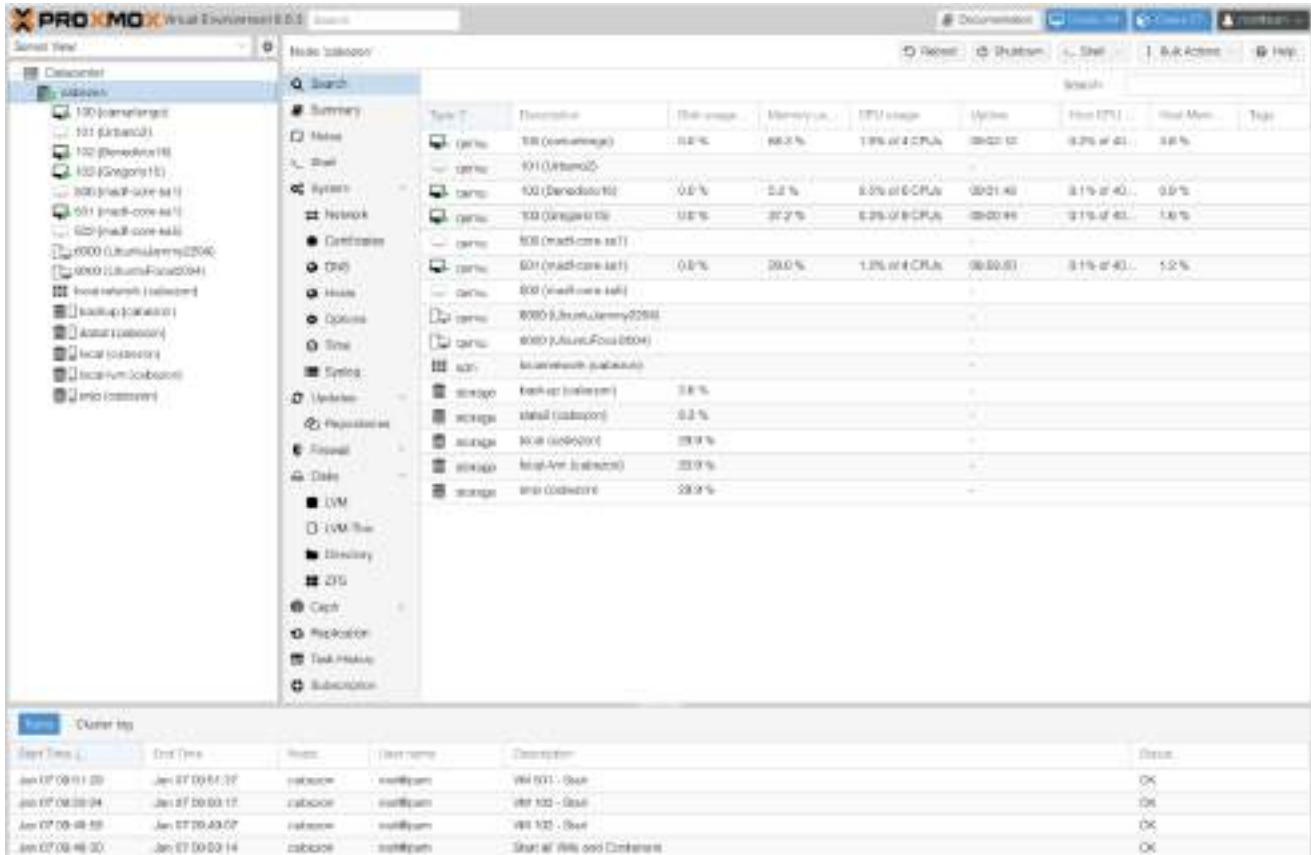
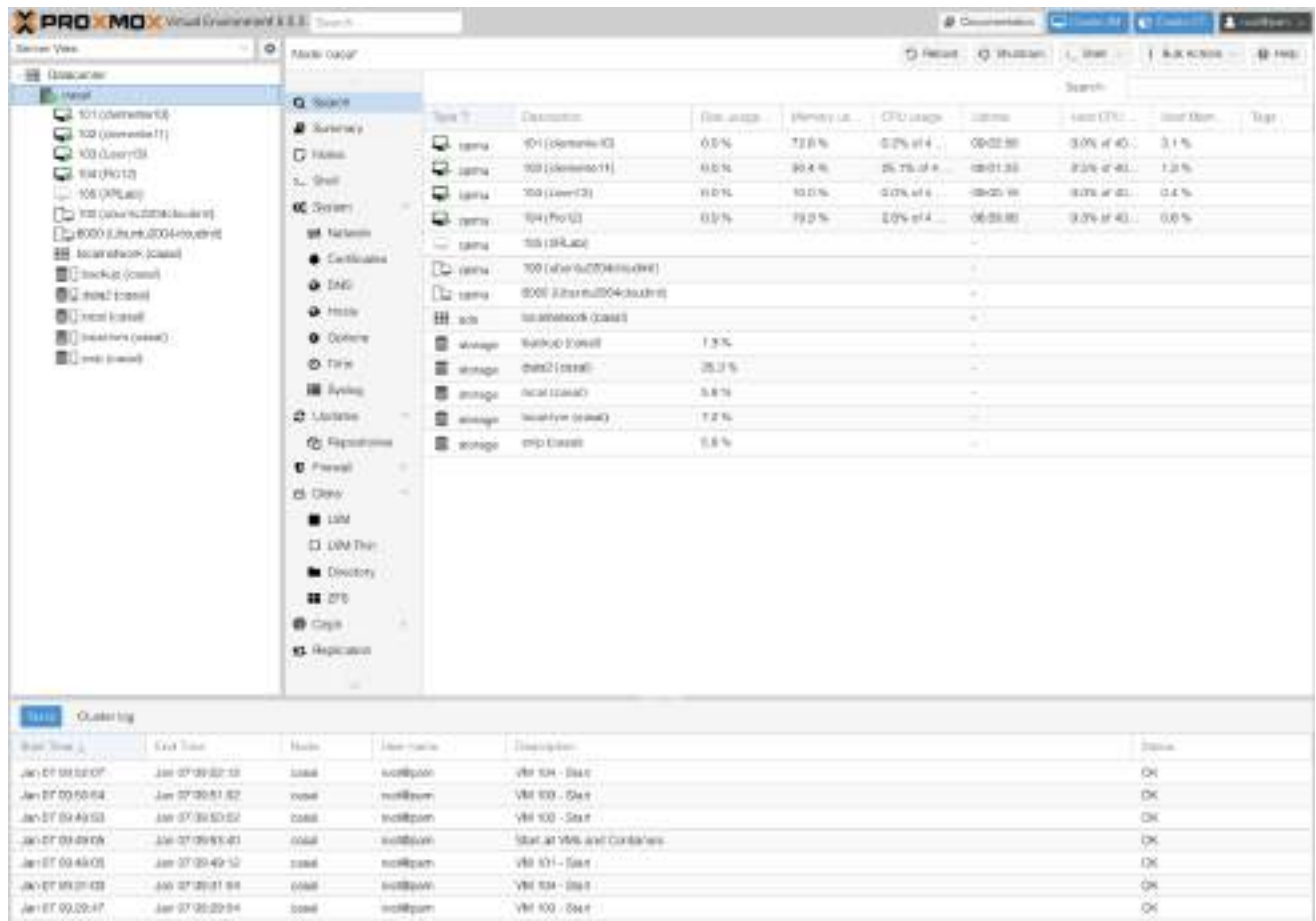


Fig. 4.2. Consola de Proxmox para el servidor OpenEdge MEC *Cabezon*

De forma complementaria, el servidor Casals concentra las capacidades de computación acelerada mediante GPU, necesarias para el offloading de algoritmos de IA. Desde su consola Proxmox es posible observar la distribución de máquinas virtuales orientadas a la ejecución de servicios de IA y backends de aplicación, así como el reparto de recursos de CPU, memoria y dispositivos PCIe asociados a las GPUs.



Type	Description	Disk usage	Memory us...	CPU usage	Status	Last CPU...	Last Mem...	Tags
lxc	VM1 (OpenEdge)	0.0 %	73.0 %	0.0% of 4	OK	0.0% of 40...	0.1 %	
lxc	VM2 (OpenEdge)	0.0 %	80.4 %	0.0% of 4	OK	0.0% of 40...	0.2 %	
lxc	VM3 (OpenEdge)	0.0 %	80.0 %	0.0% of 4	OK	0.0% of 40...	0.4 %	
lxc	VM4 (Pro 12)	0.0 %	79.0 %	0.0% of 4	OK	0.0% of 40...	0.0 %	
lxc	VM5 (OpenEdge)	-	-	-	-	-	-	
lxc	VM6 (OpenEdge)	-	-	-	-	-	-	
lxc	VM7 (OpenEdge)	-	-	-	-	-	-	
lxc	VM8 (OpenEdge)	-	-	-	-	-	-	
lxc	VM9 (OpenEdge)	-	-	-	-	-	-	
lxc	VM10 (OpenEdge)	-	-	-	-	-	-	
lxc	VM11 (OpenEdge)	-	-	-	-	-	-	
lxc	VM12 (OpenEdge)	-	-	-	-	-	-	
lxc	VM13 (OpenEdge)	-	-	-	-	-	-	
lxc	VM14 (OpenEdge)	-	-	-	-	-	-	
lxc	VM15 (OpenEdge)	-	-	-	-	-	-	
lxc	VM16 (OpenEdge)	-	-	-	-	-	-	
lxc	VM17 (OpenEdge)	-	-	-	-	-	-	
lxc	VM18 (OpenEdge)	-	-	-	-	-	-	
lxc	VM19 (OpenEdge)	-	-	-	-	-	-	
lxc	VM20 (OpenEdge)	-	-	-	-	-	-	
lxc	VM21 (OpenEdge)	-	-	-	-	-	-	
lxc	VM22 (OpenEdge)	-	-	-	-	-	-	
lxc	VM23 (OpenEdge)	-	-	-	-	-	-	
lxc	VM24 (OpenEdge)	-	-	-	-	-	-	
lxc	VM25 (OpenEdge)	-	-	-	-	-	-	
lxc	VM26 (OpenEdge)	-	-	-	-	-	-	
lxc	VM27 (OpenEdge)	-	-	-	-	-	-	
lxc	VM28 (OpenEdge)	-	-	-	-	-	-	
lxc	VM29 (OpenEdge)	-	-	-	-	-	-	
lxc	VM30 (OpenEdge)	-	-	-	-	-	-	
lxc	VM31 (OpenEdge)	-	-	-	-	-	-	
lxc	VM32 (OpenEdge)	-	-	-	-	-	-	
lxc	VM33 (OpenEdge)	-	-	-	-	-	-	
lxc	VM34 (OpenEdge)	-	-	-	-	-	-	
lxc	VM35 (OpenEdge)	-	-	-	-	-	-	
lxc	VM36 (OpenEdge)	-	-	-	-	-	-	
lxc	VM37 (OpenEdge)	-	-	-	-	-	-	
lxc	VM38 (OpenEdge)	-	-	-	-	-	-	
lxc	VM39 (OpenEdge)	-	-	-	-	-	-	
lxc	VM40 (OpenEdge)	-	-	-	-	-	-	
lxc	VM41 (OpenEdge)	-	-	-	-	-	-	
lxc	VM42 (OpenEdge)	-	-	-	-	-	-	
lxc	VM43 (OpenEdge)	-	-	-	-	-	-	
lxc	VM44 (OpenEdge)	-	-	-	-	-	-	
lxc	VM45 (OpenEdge)	-	-	-	-	-	-	
lxc	VM46 (OpenEdge)	-	-	-	-	-	-	
lxc	VM47 (OpenEdge)	-	-	-	-	-	-	
lxc	VM48 (OpenEdge)	-	-	-	-	-	-	
lxc	VM49 (OpenEdge)	-	-	-	-	-	-	
lxc	VM50 (OpenEdge)	-	-	-	-	-	-	
lxc	VM51 (OpenEdge)	-	-	-	-	-	-	
lxc	VM52 (OpenEdge)	-	-	-	-	-	-	
lxc	VM53 (OpenEdge)	-	-	-	-	-	-	
lxc	VM54 (OpenEdge)	-	-	-	-	-	-	
lxc	VM55 (OpenEdge)	-	-	-	-	-	-	
lxc	VM56 (OpenEdge)	-	-	-	-	-	-	
lxc	VM57 (OpenEdge)	-	-	-	-	-	-	
lxc	VM58 (OpenEdge)	-	-	-	-	-	-	
lxc	VM59 (OpenEdge)	-	-	-	-	-	-	
lxc	VM60 (OpenEdge)	-	-	-	-	-	-	
lxc	VM61 (OpenEdge)	-	-	-	-	-	-	
lxc	VM62 (OpenEdge)	-	-	-	-	-	-	
lxc	VM63 (OpenEdge)	-	-	-	-	-	-	
lxc	VM64 (OpenEdge)	-	-	-	-	-	-	
lxc	VM65 (OpenEdge)	-	-	-	-	-	-	
lxc	VM66 (OpenEdge)	-	-	-	-	-	-	
lxc	VM67 (OpenEdge)	-	-	-	-	-	-	
lxc	VM68 (OpenEdge)	-	-	-	-	-	-	
lxc	VM69 (OpenEdge)	-	-	-	-	-	-	
lxc	VM70 (OpenEdge)	-	-	-	-	-	-	
lxc	VM71 (OpenEdge)	-	-	-	-	-	-	
lxc	VM72 (OpenEdge)	-	-	-	-	-	-	
lxc	VM73 (OpenEdge)	-	-	-	-	-	-	
lxc	VM74 (OpenEdge)	-	-	-	-	-	-	
lxc	VM75 (OpenEdge)	-	-	-	-	-	-	
lxc	VM76 (OpenEdge)	-	-	-	-	-	-	
lxc	VM77 (OpenEdge)	-	-	-	-	-	-	
lxc	VM78 (OpenEdge)	-	-	-	-	-	-	
lxc	VM79 (OpenEdge)	-	-	-	-	-	-	
lxc	VM80 (OpenEdge)	-	-	-	-	-	-	
lxc	VM81 (OpenEdge)	-	-	-	-	-	-	
lxc	VM82 (OpenEdge)	-	-	-	-	-	-	
lxc	VM83 (OpenEdge)	-	-	-	-	-	-	
lxc	VM84 (OpenEdge)	-	-	-	-	-	-	
lxc	VM85 (OpenEdge)	-	-	-	-	-	-	
lxc	VM86 (OpenEdge)	-	-	-	-	-	-	
lxc	VM87 (OpenEdge)	-	-	-	-	-	-	
lxc	VM88 (OpenEdge)	-	-	-	-	-	-	
lxc	VM89 (OpenEdge)	-	-	-	-	-	-	
lxc	VM90 (OpenEdge)	-	-	-	-	-	-	
lxc	VM91 (OpenEdge)	-	-	-	-	-	-	
lxc	VM92 (OpenEdge)	-	-	-	-	-	-	
lxc	VM93 (OpenEdge)	-	-	-	-	-	-	
lxc	VM94 (OpenEdge)	-	-	-	-	-	-	
lxc	VM95 (OpenEdge)	-	-	-	-	-	-	
lxc	VM96 (OpenEdge)	-	-	-	-	-	-	
lxc	VM97 (OpenEdge)	-	-	-	-	-	-	
lxc	VM98 (OpenEdge)	-	-	-	-	-	-	
lxc	VM99 (OpenEdge)	-	-	-	-	-	-	
lxc	VM100 (OpenEdge)	-	-	-	-	-	-	

Fig. 4.3. Consola de Proxmox para el servidor OpenEdge MEC *Casals*

El entorno de virtualización permite, además, la monitorización detallada del estado de cada VNF. A través de la consola se visualizan métricas de uso de CPU, memoria y almacenamiento, lo que facilita la detección temprana de posibles cuellos de botella y la validación del correcto funcionamiento de los distintos componentes desplegados.

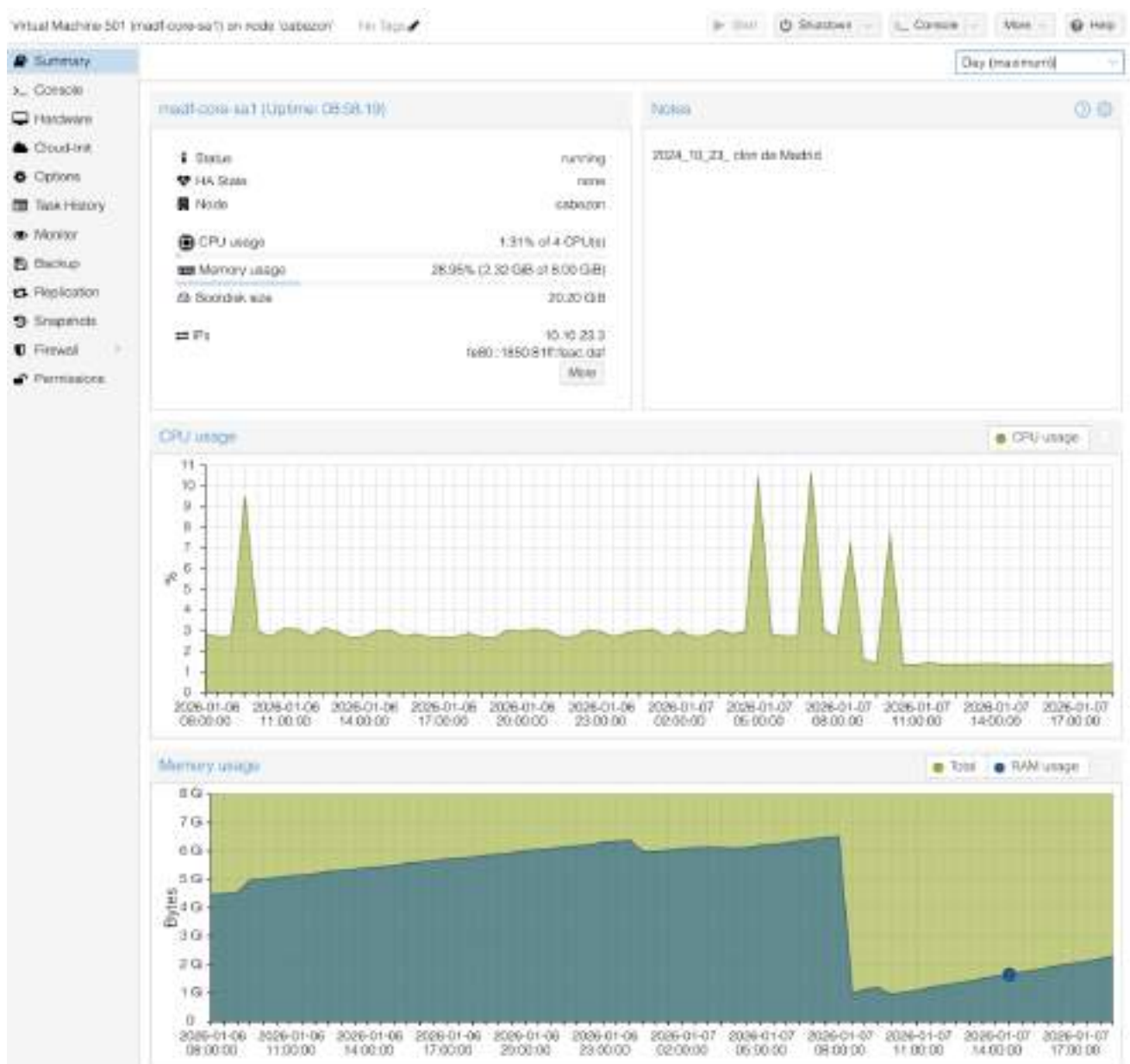


Fig. 4.4. Monitorización de una de las VNFs de *Casals* (core 5G) desde la consola Proxmox. Asimismo, se muestra el seguimiento de VNFs desplegadas en el servidor Cabezon, confirmando que los servicios de soporte del MEC operan de forma estable y continua durante las pruebas funcionales.

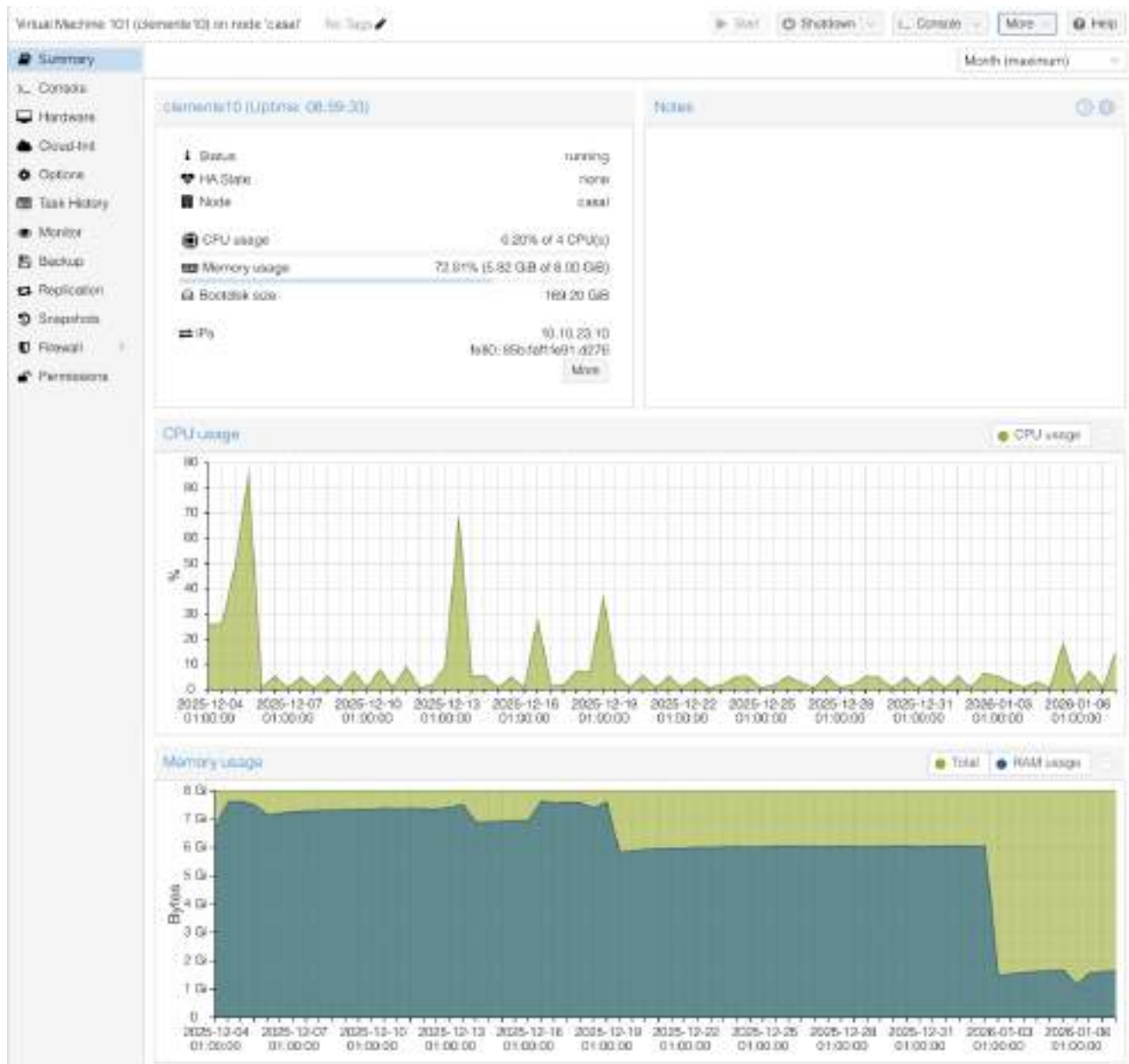


Fig. 4.5. Monitorización de una de las VNFs de *Cabezon (Clemente10)* desde la consola Proxmox

4.1.3 VNF de offloading de algoritmos de IA para XR

Las funciones de offloading de algoritmos de inteligencia artificial para los casos de uso XR se implementan mediante VNFs específicas desplegadas sobre el servidor Casals. Estas VNFs están diseñadas para ejecutar algoritmos de segmentación y procesamiento de imágenes de forma acelerada, aprovechando las GPUs disponibles en el nodo MEC.

El software de estos servicios se encapsula en contenedores Docker, lo que permite desplegar distintas versiones de algoritmos y facilitar su actualización, activación o desactivación de manera dinámica. En la siguiente figura se muestran las imágenes Docker instaladas en una de las VNFs de offloading, correspondientes a los distintos algoritmos de IA empleados en el proyecto.

```
xruser@clemente10:~$ docker images
```

REPOSITORY	TAG	IMAGE ID	CREATED	SIZE
xrlab/ego-segmentation	latest	fe0df4148fa4	4 weeks ago	17.3GB
<none>	<none>	f40d331251b9	4 weeks ago	17.3GB
<none>	<none>	def3d3ce9079	4 weeks ago	17.3GB
<none>	<none>	80973a0b58aa	4 weeks ago	17GB
xrlab/emotions	latest	9aa8b9547e24	6 weeks ago	2.52GB
xrlab/adapchroma	latest	4cec003de663	6 weeks ago	12.7GB
gryphon	1.0	8b6c6dbff0df	2 months ago	1.16GB
polyp	2.3dev	f36dd6c91a78	2 months ago	1.51GB
yolothon	latest	f72934d8b7da	13 months ago	32.5GB
hello-world	latest	d2c94e258dcb	2 years ago	13.3kB

Fig. 4.6. Imágenes docker instaladas en *Clemente10* para ejecutar los distintos algoritmos de offloading de IA

La correcta utilización de los recursos de aceleración se valida mediante herramientas de monitorización específicas de GPU. En este contexto, se emplea la utilidad nvidia-smi para observar en tiempo real la carga de las GPUs, el consumo de memoria gráfica y los procesos asociados a la ejecución de los algoritmos de IA.

```
xruser@clemente10:~$ nvidia-smi
```

Wed Jan 7 17:12:33 2026

NVIDIA-SMI 470.256.02 Driver Version: 470.256.02 CUDA Version: 11.4									
GPU	Name	Persistence-M	Bus-Id	Disp.A	Volatile	Uncorr. ECC			
Fan	Temp	Perf	Pwr:Usage/Cap	Memory-Usage	GPU-Util	Compute M.			
MIG M.									
0	Tesla T4	Off	00000000:00:18.0	Off	11%	0			
N/A	57C	P0	44W / 70W	2173MiB / 15109MiB		Default			
							N/A		

Processes:							
GPU	GI	CI	PID	Type	Process name	GPU Memory Usage	
ID	ID	ID					
0	N/A	N/A	11167	C	python3	2170MiB	

Fig. 4.7. Monitorización de la GPU con Nvidia-SMI

De forma complementaria, se supervisa el consumo de CPU y la actividad de los procesos dentro de la VNF mediante herramientas estándar de sistema. Esto permite verificar que la

carga computacional se distribuye adecuadamente entre CPU y GPU, y que no se producen situaciones de saturación durante la ejecución de los servicios de offloading.

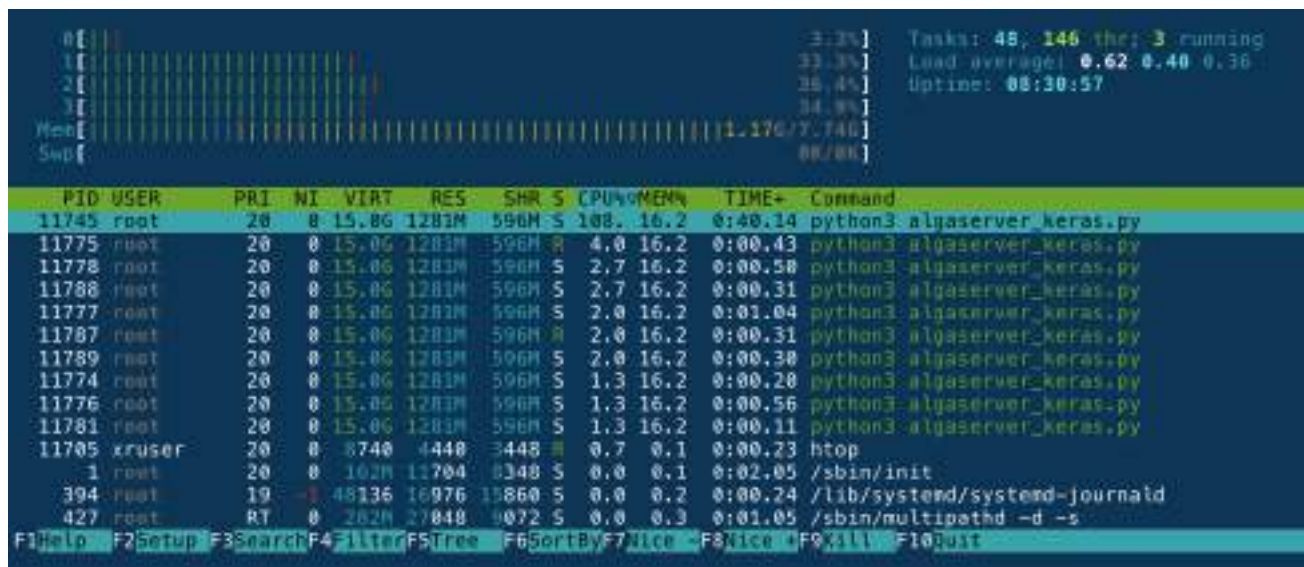


Fig. 4.8. Monitorización de la CPU con htop

Finalmente, se muestra la ejecución simultánea de contenedores en una segunda VNF de offloading, lo que evidencia la capacidad del sistema para escalar horizontalmente la ejecución de algoritmos de IA y dar soporte concurrente a distintos casos de uso.



Fig. 4.9. Imágenes docker ejecutando en *Clemente11* para dar servicio a los casos de uso

4.1.4 VNF de backend de telepresencia

El backend del sistema de telepresencia se despliega como una VNF independiente en el nodo MEC, encargada de gestionar tanto el plano de datos como el plano de control del servicio. Esta VNF da soporte al caso de uso de teleportación virtual, actuando como punto de terminación de los flujos de vídeo y de señalización asociados.

Al igual que en el resto de servicios del sistema edge, el backend de telepresencia se basa en una arquitectura modular soportada por contenedores Docker. Esta aproximación permite separar claramente los distintos componentes funcionales y facilita su mantenimiento y evolución.

En la siguiente figura se muestran los contenedores activos en la VNF correspondiente, evidenciando la correcta ejecución de los servicios necesarios para el funcionamiento del sistema de telepresencia durante las pruebas funcionales.

```

xcruser@Pio12:~$ docker ps
CONTAINER ID   IMAGE          COMMAND                  CREATED        STATUS        PORTS        NAMES
21a8ca6d2e5c   pulpo:1.8     "/docker-entrypoint..." 5 months ago   Up 9 hours   -            pulpo
8d733e76703e   polyp:2.2     "/docker-entrypoint..." 5 months ago   Up 9 hours   -            polyp
619eec01f005   gryphon:1.8   "/entrypoint.sh tele..." 18 months ago  Up 9 hours   -            gryphon

```

Fig. 4.10. Imágenes docker ejecutando en *Pio12* (backend de telepresencia)

El contenedor *pulpo* contiene el plano de gestión del backend, el contenedor *polyp* el plano de datos, y el contenedor *gryphon* el sistema de monitorización.

4.1.5 VNF de backend de teleformación

El backend de teleformación y terapia se implementa mediante una VNF específica encargada de soportar la lógica de aplicación de los casos de uso asociados a formación y sesiones terapéuticas. Esta VNF proporciona los servicios de control, configuración y comunicación necesarios para la interacción entre los usuarios finales y el sistema edge.

La arquitectura del backend se basa en un servidor de aplicaciones desarrollado sobre Node.js, utilizando Express para la provisión de interfaces web y Socket.IO para la mensajería en tiempo real. Esta combinación permite una gestión flexible de las sesiones y una comunicación eficiente con el resto de componentes del sistema.

La figura correspondiente muestra los contenedores Docker activos en esta VNF, confirmando el despliegue correcto de los servicios de backend y su disponibilidad durante las pruebas.

```

xcruser@Leon13:~$ docker ps
CONTAINER ID   IMAGE                  COMMAND                  CREATED        STATUS        PORTS        NAMES
3d5b386de158   arlab/drcommands-server "/docker-entrypoint.s..." 11 months ago  Up 9 hours   0.0.0.0:3001->3001/tcp, ...  drcommands-server
364e2a28b778   arlab/node-controller "/docker-entrypoint.s..." 14 months ago  Up 9 hours   -            comdia

```

Fig. 4.11. Imágenes docker ejecutando en *Leon13* (backend de teleformación / terapia)

Asimismo, una segunda VNF instala un CMS basado en MediaCMS para almacenar los contenidos inmersivos de dichos casos de uso.

```

xcruser@Gregorio15:~$ docker ps
CONTAINER ID   IMAGE                  COMMAND                  CREATED        STATUS        PORTS        NAMES
0ae11514e01b   charmell/soft-server:latest "/usr/local/bin/soft..." 4 weeks ago   Up 9 hours   0.0.0.0:9418->9418/tcp, ...  soft-server
8.8.8.0:23231->23231->23231/tcp, ... 23231-23231->23231-23231/tcp
f1ae6df1230a   mediacms/mediacms:latest "/deploy/docker/ent..." 7 months ago  Up 9 hours   80/tcp, 5008/tcp      mediacms-gallery-worker-1
Baeec9f8eadd   mediacms/mediacms:latest "/deploy/docker/ent..." 7 months ago  Up 9 hours   0.0.0.0:80->80/tcp, ...  mediacms-ent-1
c796e781e01e   mediacms/mediacms:latest "/deploy/docker/ent..." 7 months ago  Up 9 hours   80/tcp, 5008/tcp      mediacms-gallery-worker-2
fa44fc0e77fe   gowtgrek:15.2-alpine "/docker-entrypoint.s..." 14 months ago  Up 9 hours (healthy)  5412/tcp              mediacms-db-1
438ee7950e13   radia:alpine         "/docker-entrypoint.s..." 14 months ago  Up 9 hours (healthy)  6379/tcp              mediacms-media-2

```

Fig. 4.12. Imágenes docker ejecutando en *Gregorio15* (CMS para teleformación / terapia)

4.1.6 VNF de monitorización y gestión de KPIs

Con el objetivo de disponer de una visión global del estado del sistema y de sus prestaciones, se ha desplegado una VNF dedicada a la monitorización y gestión de

indicadores clave de rendimiento (KPIs). Esta VNF constituye un elemento fundamental para la validación experimental del sistema de edge y para el análisis posterior de resultados.

La monitorización se basa en el stack TIG, compuesto por Telegraf para la recolección de métricas, InfluxDB para su almacenamiento y Grafana para la visualización. Este conjunto permite integrar métricas procedentes de la red, de las máquinas virtuales, de los contenedores y de las aplicaciones.

Los servicios que conforman este stack se gestionan mediante systemd, asegurando su arranque automático y su operación continua. La siguiente figura muestra el estado de dichos servicios en la VNF de monitorización, confirmando su correcto despliegue y funcionamiento.

```
xruser@Benedicto16:~$ systemctl status telegraf.service
● telegraf.service - Telegraf
   Loaded: loaded (/lib/systemd/system/telegraf.service; enabled; vendor preset: enabled)
   Active: active (running) since Wed 2026-01-07 08:49:30 UTC; 8h ago
     Docs: https://github.com/influxdata/telegraf
   Main PID: 696 (telegraf)
    Tasks: 13 (limit: 38464)
   Memory: 100.7M
      CPU: 3min 5.235s
   CGroup: /system.slice/telegraf.service
           └─696 /usr/bin/telegraf -config /etc/telegraf/telegraf.conf --config-directory /etc/telegraf/telegraf.d

xruser@Benedicto16:~$ systemctl status influxdb.service
● influxdb.service - InfluxDB is an open-source, distributed, time series database
   Loaded: loaded (/lib/systemd/system/influxdb.service; enabled; vendor preset: enabled)
   Active: active (running) since Wed 2026-01-07 08:49:31 UTC; 8h ago
     Docs: https://docs.influxdata.com/influxdb/
  Process: 682 ExecStart=/usr/lib/influxdb/scripts/influxd-systemd-start.sh (code=exited; status=0/SUCCESS)
   Main PID: 705 (influxd)
    Tasks: 25 (limit: 38464)
   Memory: 493.6M
      CPU: 3min 17.133s
   CGroup: /system.slice/influxdb.service
           └─705 /usr/bin/influxd

xruser@Benedicto16:~$ systemctl status grafana-server.service
● grafana-server.service - Grafana instance
   Loaded: loaded (/lib/systemd/system/grafana-server.service; enabled; vendor preset: enabled)
   Active: active (running) since Wed 2026-01-07 08:49:31 UTC; 8h ago
     Docs: http://docs.grafana.org
   Main PID: 848 (grafana)
    Tasks: 16 (limit: 38464)
   Memory: 355.4M
      CPU: 1min 57.957s
   CGroup: /system.slice/grafana-server.service
           └─848 /usr/share/grafana/bin/grafana server --config=/etc/grafana/grafana.ini --pidfile=/run/graf
```

4.2

Fig. 4.13. Servicios del Stack TIG configurados en *Benedicto16* mediante systemd

Pruebas de carga

En este apartado se presentan las pruebas de carga realizadas sobre el sistema de edge con el objetivo de caracterizar el rendimiento computacional de las funciones de offloading de algoritmos de inteligencia artificial desplegadas en el nodo MEC. A diferencia de los apartados anteriores, centrados en la validación del comportamiento de la red 5G, en este caso el foco se sitúa en la capacidad de cómputo del sistema, y en particular en el tiempo de ejecución de los algoritmos de IA bajo distintas configuraciones.

Las pruebas se estructuran en dos escenarios, alineados con los casos de uso del proyecto: el procesamiento delegado de Realidad Mixta y la Regulación Emocional. Para ambos escenarios se analizan métricas relacionadas con el tiempo de inferencia de los algoritmos, a partir de las cuales se derivan indicadores como la tasa de cuadros por segundo (FPS) y el consumo estimado de recursos.

De forma previa, la siguiente tabla resume los valores objetivo y órdenes de magnitud esperados para cada escenario, en términos de tasa de cuadros, resolución, tasa binaria de entrada y salida, tiempo de inferencia y uso de memoria GPU, sirviendo como referencia para la interpretación de los resultados experimentales.

Tabla 20. Requisitos y valores de referencia para las pruebas de carga de los escenarios de computación en el edge

	Computación Delegada en MR	Regulación Emocional
Tasa de cuadro	15 - 60 fps	10 muestras / s
Resolución de cuadro	1280 x 480	N/A
Tasa binaria E/S	~100 Mbps	~100 kbps
Tiempo de inferencia	10-20 ms	0.1 – 1 s
Memoria GPU estimada	< 4GB	< 4GB

4.2.1 Escenario 2: Procesamiento delegado de Realidad Mixta (*Mixed Reality computation offloading*)

En este escenario se evalúa el rendimiento del sistema de edge para la ejecución de algoritmos de segmentación y procesamiento de imagen asociados al caso de uso de Realidad Mixta. El objetivo principal es caracterizar el tiempo de inferencia de distintos algoritmos de IA cuando se ejecutan en el nodo MEC, considerando diferentes resoluciones de entrada.

Configuración y procedimiento de la prueba

Las pruebas se han realizado ejecutando los algoritmos de forma individual, sin concurrencia entre ellos, con el fin de aislar su comportamiento y evitar interferencias. Cada prueba tiene una duración aproximada de tres minutos, durante los cuales se registra el tiempo de ejecución de cada ciclo de inferencia.

Todos los algoritmos se ejecutan en la VNF *Clemente10*, desplegada como una máquina virtual con acceso en passthrough a una GPU Nvidia Tesla T4. Los recursos asignados a esta VNF son:

- CPU: 4 vCPU de arquitectura x86_64

- Memoria RAM: 8 GB
- GPU: Nvidia Tesla T4 (passthrough)

El sistema empleado para la ejecución de las pruebas es el mismo sistema SBSP utilizado en las pruebas de red de este escenario, si bien en este caso no se analizan métricas de red. En su lugar, se mide el tiempo de ejecución de cada inferencia, que constituye la métrica primaria de estas pruebas.

A partir del tiempo de inferencia registrado, se calculan el resto de indicadores presentados en este apartado. En particular, la tasa de cuadros por segundo (FPS) se deriva considerando un overhead constante por frame, común a todos los algoritmos, lo que permite comparar de forma homogénea su rendimiento relativo.

Las métricas se recogen y visualizan mediante paneles de Grafana, configurados para mostrar la evolución temporal del tiempo de inferencia y sus estadísticos básicos (valor medio, máximo, mínimo y desviación estándar).

Algoritmos evaluados

Para este escenario se han evaluado los siguientes algoritmos de procesamiento y segmentación, de acuerdo con la configuración definida en el proyecto:

- **Ego Body segmentation**
 - Modelo Keras.hdf5 (Thundernet)
- **Algoritmo adaptativo**
 - Segmentación basada en Chroma con umbrales adaptativos derivados de Thundernet
- **Segmentación de objetos**
 - YOLO-cutlery, en distintas combinaciones con Thundernet, EgoHOS y Chroma

Las pruebas se han realizado para dos resoluciones de entrada, representativas de los modos de operación del sistema: 1280×480 y 2540×960 píxeles.

4.2.1.1 Thundernet

En primer lugar se analiza el comportamiento del algoritmo Thundernet, empleado para la segmentación del cuerpo del usuario. La siguiente figura muestra la evolución temporal del tiempo de inferencia para la resolución de 1280×480 píxeles, junto con los valores estadísticos asociados.

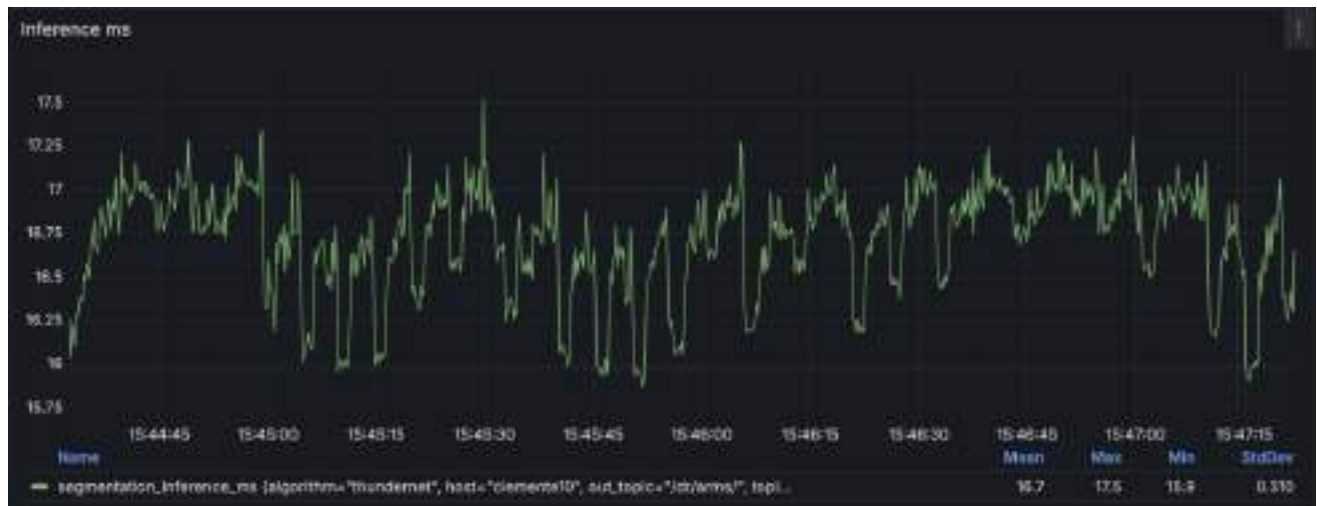


Fig. 4.14. Tiempo de inferencia para el algoritmo Thundernet con resolución 1280x480

Se observa un comportamiento estable, con valores de inferencia concentrados en torno a su valor medio y una desviación estándar reducida. Para esta resolución, el tiempo medio de inferencia se sitúa en el orden de decenas de milisegundos.

A continuación se presenta el resultado correspondiente a la resolución de 2540x960 píxeles.

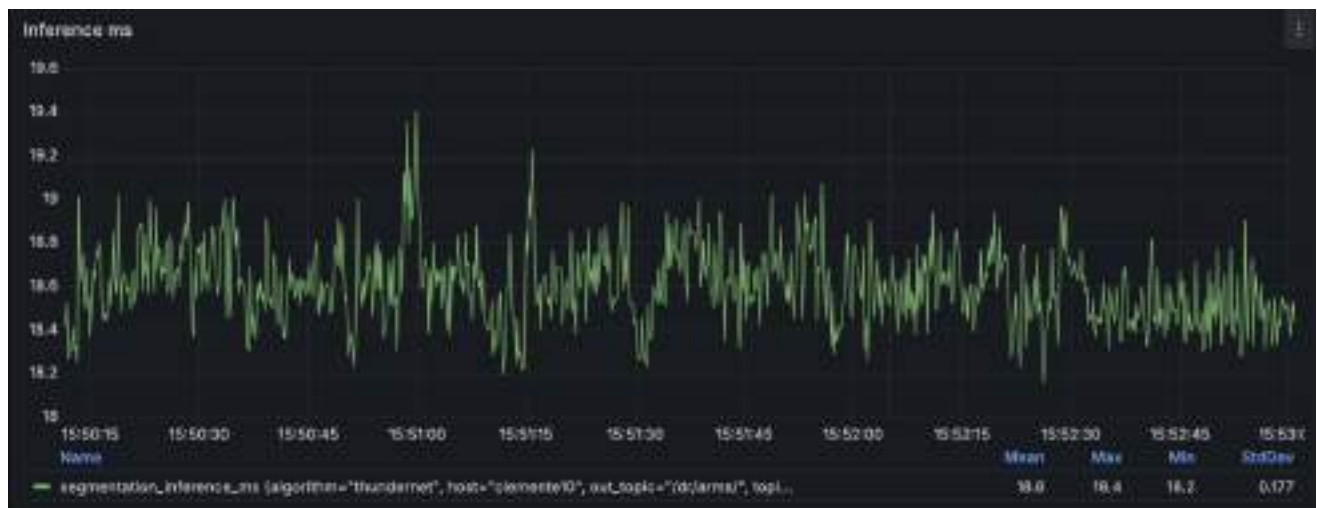
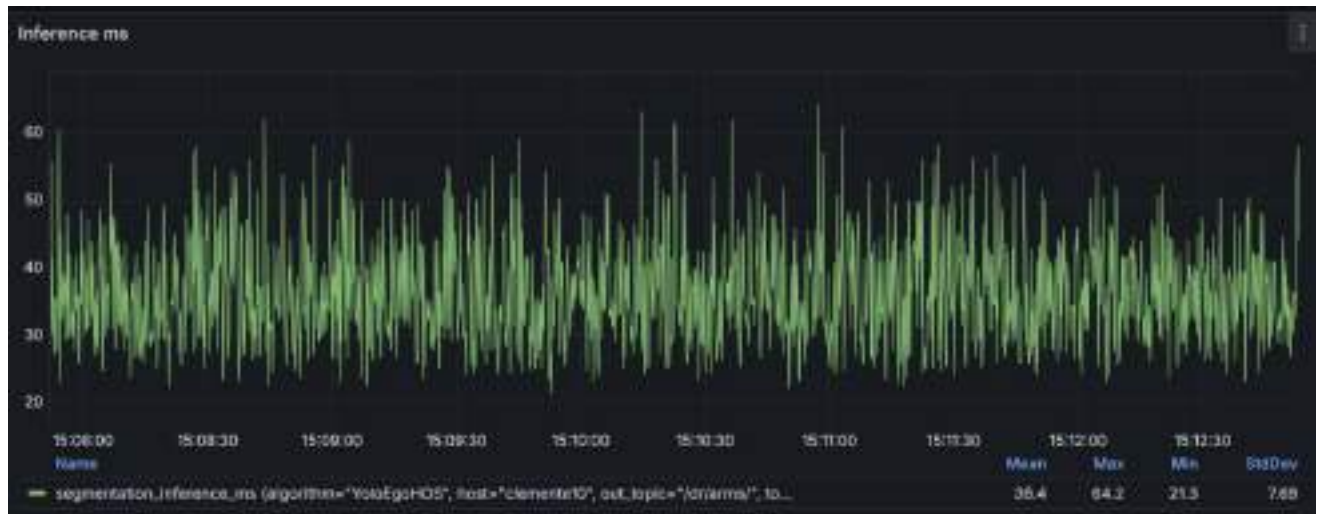


Fig. 4.15. Tiempo de inferencia para el algoritmo Thundernet con resolución 2540x960

Como era esperable, el incremento de resolución conlleva un aumento moderado del tiempo de inferencia, manteniéndose no obstante una variabilidad baja y un comportamiento consistente a lo largo de la prueba.

4.2.1.2 Yolo + EgoHos

En este apartado se presentan los resultados del algoritmo de segmentación de objetos basado en YOLO combinado con EgoHOS. Para la resolución de 1280×480 píxeles, la siguiente figura recoge la evolución del tiempo de inferencia durante la prueba.



4.2.1.3 Fig. 4.16. Tiempo de inferencia para el algoritmo Yolo + EgoHos con resolución 1280x480

En comparación con Thundernet, se observa un aumento significativo tanto del valor medio como de la dispersión de los tiempos de inferencia, reflejando una mayor complejidad computacional del algoritmo combinado.

La figura siguiente muestra los resultados obtenidos para la resolución de 2540×960 píxeles.

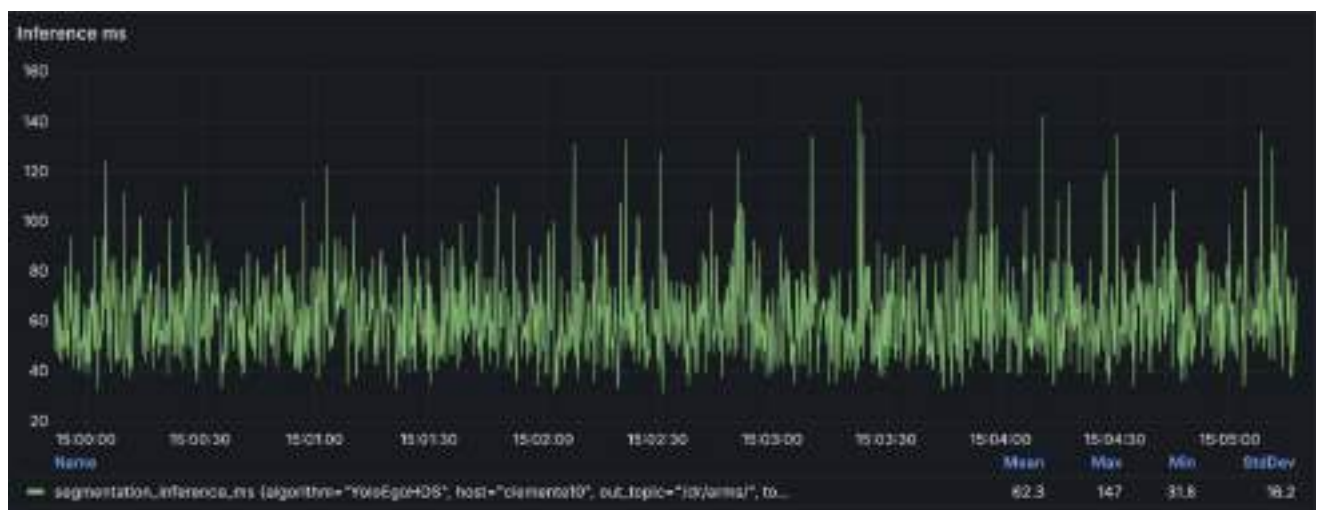


Fig. 4.17. Tiempo de inferencia para el algoritmo Yolo + EgoHos con resolución 2540 x 960

En este caso, el aumento de resolución se traduce en un incremento notable del tiempo de inferencia medio y de la variabilidad, alcanzándose valores máximos sensiblemente superiores a los observados en resoluciones más bajas.

4.2.1.4 Yolo + Thundernet

Se evalúa a continuación la combinación de YOLO con Thundernet. La siguiente figura corresponde a la resolución de 1280×480 píxeles.

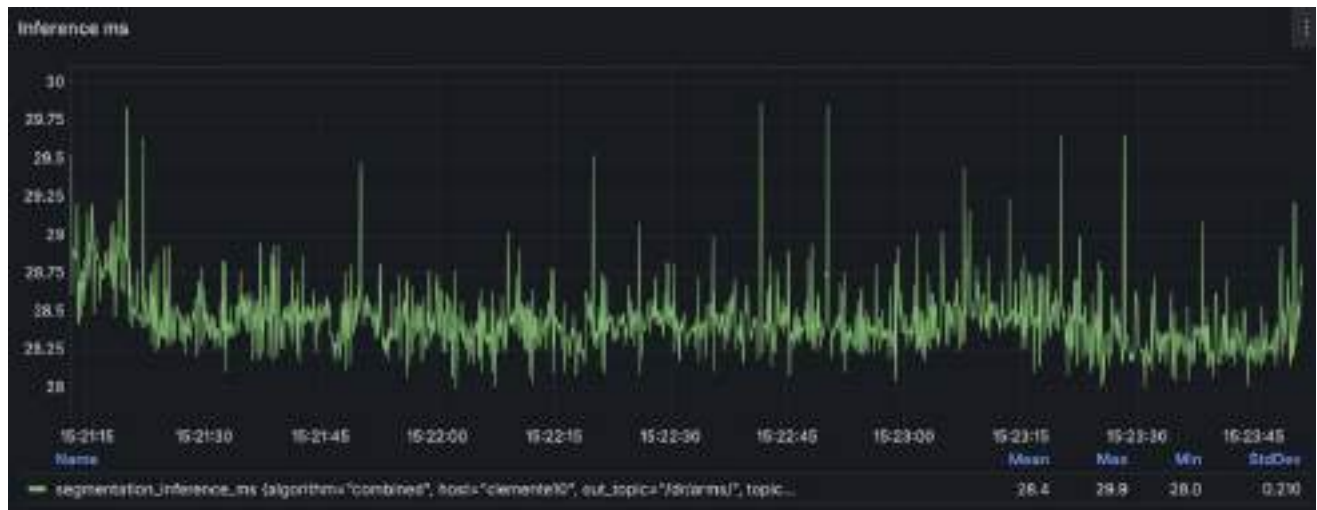


Fig. 4.18. Tiempo de inferencia para el algoritmo Yolo +Thundernet con resolución 1280x480

Los resultados muestran un comportamiento intermedio entre los algoritmos evaluados previamente, con tiempos de inferencia superiores a Thundernet pero inferiores a la combinación Yolo + EgoHOS, y una variabilidad contenida.

Para la resolución de 2540×960 píxeles se obtienen los resultados mostrados a continuación.

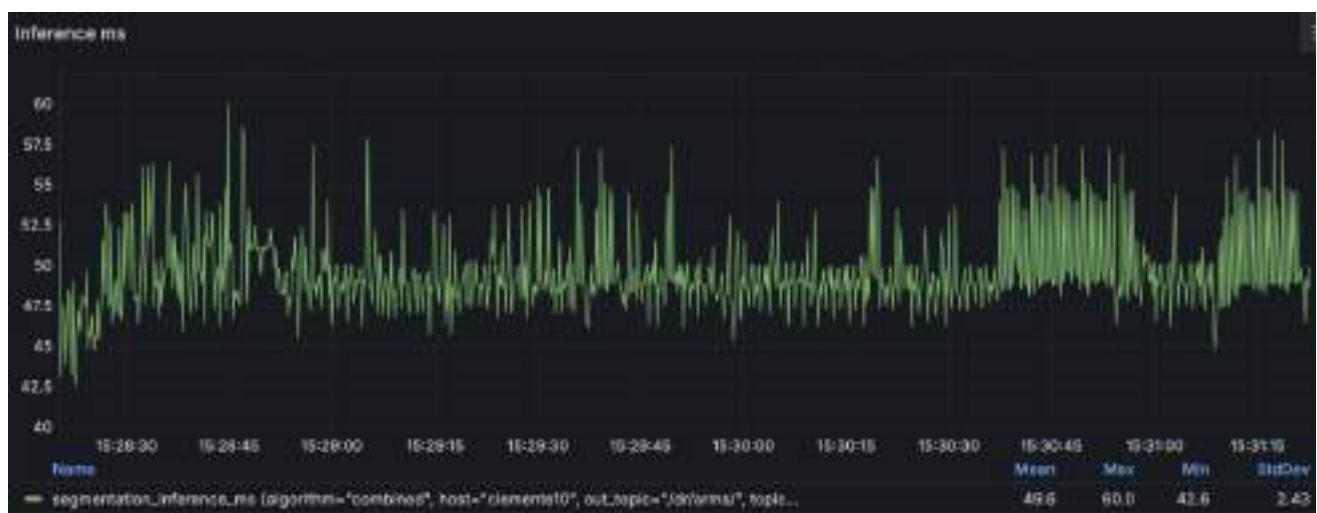


Fig. 4.19. Tiempo de inferencia para el algoritmo Yolo + Thundernet con resolución 2540x960

El incremento de resolución provoca un aumento apreciable del tiempo de inferencia, si bien el algoritmo mantiene un comportamiento relativamente estable durante toda la ejecución de la prueba.

4.2.1.5 Yolo

En este subapartado se presentan los resultados correspondientes al algoritmo YOLO ejecutado de forma independiente. La siguiente figura muestra el tiempo de inferencia para la resolución de 1280x480 píxeles.



Fig. 4.20. Tiempo de inferencia para el algoritmo Yolo con resolución 1280x480

Los valores obtenidos reflejan tiempos de inferencia reducidos y una baja variabilidad, lo que se traduce en tasas de procesamiento elevadas.

Para la resolución de 2540x960 píxeles, los resultados se muestran en la figura siguiente.



Fig. 4.21. Tiempo de inferencia para el algoritmo Thudernet con resolución 2540x960

Aunque el aumento de resolución incrementa el tiempo de inferencia, los valores se mantienen en rangos compatibles con los requisitos del escenario de Realidad Mixta.

4.2.1.6 Yolo + Chroma

A continuación se evalúa la combinación de YOLO con segmentación basada en Chroma.

Para la resolución de 1280x480 píxeles, la evolución del tiempo de inferencia se presenta en la siguiente figura.



Fig. 4.22. Tiempo de inferencia para el algoritmo Yolo + Chroma con resolución 1280x480

El comportamiento observado es estable, con tiempos de inferencia moderados y una dispersión limitada.

La figura siguiente muestra los resultados correspondientes a la resolución de 2540x960 píxeles.



Fig. 4.23. Tiempo de inferencia para el algoritmo Yolo + Chroma con resolución 2540x960

En este caso se aprecia un incremento del tiempo de inferencia asociado al aumento de resolución, manteniéndose no obstante un comportamiento regular durante la prueba.

4.2.1.7 Chroma

En este apartado se presentan los resultados del algoritmo Chroma ejecutado de forma aislada. La siguiente figura corresponde a la resolución de 1280x480 píxeles.



Fig. 4.24. Tiempo de inferencia para el algoritmo Chroma con resolución 1280x480

Los tiempos de inferencia obtenidos son significativamente inferiores a los de los algoritmos basados en redes neuronales profundas, con valores del orden de pocos milisegundos.

Para la resolución de 2540x960 píxeles, los resultados se muestran a continuación.



Fig. 4.25. Tiempo de inferencia para el algoritmo Chroma con resolución 2540x960

Aun con el incremento de resolución, el algoritmo mantiene tiempos de inferencia muy reducidos y una variabilidad mínima.

4.2.1.8 Adaptive Chroma

Finalmente, se evalúa el algoritmo de Chroma adaptativo. En la siguiente figura se muestra el tiempo de inferencia para la resolución de 1280×480 píxeles.



4.2.1.9 Fig. 4.26. Tiempo de inferencia para el algoritmo Adaptive Chroma con resolución 1280x480

Los resultados reflejan tiempos de inferencia sensiblemente superiores a los del Chroma convencional, como consecuencia de la lógica adaptativa adicional incorporada al algoritmo. La figura siguiente corresponde a la resolución de 2540×960 píxeles.



Fig. 4.27. Tiempo de inferencia para el algoritmo Adaptive Chroma con resolución 2540x960

En este caso, el aumento de resolución se traduce en un incremento notable del tiempo de inferencia medio, manteniéndose no obstante un comportamiento estable durante la prueba.

4.2.1.10 Resumen

La siguiente tabla recoge de forma consolidada los resultados obtenidos para todos los algoritmos evaluados en este escenario, incluyendo el tiempo medio de inferencia, la desviación estándar, la tasa de cuadros por segundo estimada y el consumo de memoria GPU.

Tabla 21. Resumen del Escenario 2: métricas de inferencia, tasa de cuadros y uso de memoria GPU

Algoritmo	Resolución de cuadro	Tiempo medio de inferencia (ms)	Desviación estándar (ms)	Tasa de cuadro (FPS)	Memoria (MiB)
Referencia E1.2	1280x480	10-20	-	15-60	4096
Thundernet	1280x480	16,7	0,31	59,9	2151
	2540x960	18,6	0,17	53,8	
Yolo + EgoHos	1280x480	36,4	7,69	27,5	1910
	2540x960	62,3	16,2	16,1	
Yolo + Thundernet	1280x480	28,4	0,21	35,2	14049
	2540x960	49,6	2,43	20,2	
Yolo	1280x480	11,1	0,56	90,1	407
	2540x960	15,4	1,67	64,9	
Yolo + Chroma	1280x480	14,5	0,72	69,0	407
	2540x960	26,8	0,36	37,3	
Chroma	1280x480	2,14	0,17	467,3	0
	2540x960	5,18	0,05	193,1	
Adaptive Chroma	1280x480	65,9	0,361	15,2	0
	2540x960	108	2,0	9,3	

En la tabla se incluye, a modo de referencia, una fila correspondiente a los valores objetivo definidos en el entregable E1.2. Los resultados experimentales permiten comparar de forma directa el comportamiento relativo de los distintos algoritmos y su adecuación a los requisitos del escenario de Realidad Mixta.

Cabe destacar que, en el caso de la combinación Yolo + Thundernet, se observa un consumo de memoria GPU significativamente superior al del resto de algoritmos. Este comportamiento queda identificado como un aspecto a revisar en futuras iteraciones.

4.2.2 Escenario 3: Regulación Emocional (*Emotional Regulation*)

En este escenario se analizan las prestaciones computacionales de los algoritmos empleados en el caso de uso de regulación emocional. De acuerdo con la configuración actual del sistema, se han considerado algoritmos de clasificación y análisis de señales fisiológicas ejecutados sobre datos pregrabados o sintéticos, con el objetivo de caracterizar su comportamiento en términos de tiempo de inferencia y estabilidad de ejecución.

El conjunto de algoritmos evaluados en este escenario incluye modelos de baja complejidad computacional orientados a la extracción de indicadores de estado emocional y a la clasificación de patrones asociados, siendo representativos de los componentes que se integrarían en un servicio de entrenamiento cognitivo o regulación emocional desplegado en el nodo de edge computing. Estos algoritmos se han diseñado para operar en tiempo casi real y para ser ejecutados en infraestructuras generalistas, sin requerir aceleración hardware específica.

A modo ilustrativo, la Figura 4.28 muestra el tiempo de inferencia correspondiente a un algoritmo basado en un clasificador SVM lineal, ejecutado de forma repetida sobre el conjunto de datos considerado.



Fig. 4.28. Tiempo de inferencia para el algoritmo SVM lineal para regulación emocional

Los resultados obtenidos reflejan tiempos de inferencia muy reducidos y una variabilidad limitada entre ejecuciones consecutivas, coherentes con la baja complejidad computacional de este tipo de algoritmos. Este comportamiento confirma que el escenario de regulación emocional presenta una carga de procesamiento moderada, compatible con su despliegue

en máquinas virtuales generalistas dentro del entorno MEC, y permite su ejecución simultánea junto con otros servicios XR más exigentes sin impacto apreciable en el rendimiento global de la plataforma.

4.2.3 Conclusiones

Las pruebas de carga realizadas permiten caracterizar el rendimiento computacional del sistema de edge desplegado en el proyecto, poniendo de manifiesto su capacidad para ejecutar algoritmos de IA con distintos niveles de complejidad.

En el escenario de Procesamiento delegado de Realidad Mixta, se han evaluado múltiples algoritmos de segmentación, observándose diferencias significativas en el tiempo de inferencia, la tasa de cuadros alcanzable y el consumo de memoria GPU en función del enfoque utilizado y de la resolución de entrada.

Los resultados consolidados permiten identificar configuraciones compatibles con los requisitos definidos para el escenario, así como detectar comportamientos anómalos, como el elevado consumo de memoria observado en una de las combinaciones evaluadas.

En el escenario de Regulación Emocional, los resultados disponibles muestran tiempos de inferencia reducidos para el algoritmo analizado, lo que permite simultanear una gran cantidad de usuarios concurrentes, llegado el caso.

5 Evaluación del rendimiento de la plataforma desplegada en la Fundación Juan XXIII en los distintos los casos de uso

Este capítulo presenta la evaluación del rendimiento de la plataforma desplegada en la Fundación Juan XXIII para dar soporte a los distintos casos de uso del proyecto. El análisis se orienta a valorar la capacidad del sistema para prestar estos servicios de forma estable en un entorno real, yendo un paso más allá de la validación funcional realizada en entregables previos.

Para ello, se estudia el comportamiento conjunto de la infraestructura 5G, los servicios de edge computing y los mecanismos de control de calidad en tres casos de uso representativos: terapia, telepresencia y teleformación. Los resultados obtenidos permiten discutir la viabilidad operativa de cada caso, identificando los factores limitantes y los márgenes disponibles para su despliegue futuro en condiciones de uso continuado.

5.1 Terapia conductual

5.1.1 Sistemas desplegados

El caso de uso de terapia se soporta sobre la infraestructura 5G privada desplegada en la Fundación Juan XXIII, combinando conectividad 5G Standalone en banda n40, capacidades de computación en el borde (MEC) y funciones virtualizadas específicas para el procesamiento de aplicaciones XR y de regulación emocional.

Las aplicaciones terapéuticas se ejecutan sobre VNFs desplegadas en el nodo MEC, que incluyen el backend de gestión de sesiones y contenidos, así como una VNF específica para el offloading de algoritmos de inteligencia artificial. Esta última da soporte tanto a funciones de procesamiento delegado de realidad mixta, basadas en algoritmos de segmentación semántica, como a los algoritmos empleados en el escenario de regulación emocional, caracterizados por un bajo coste computacional.

Desde el punto de vista de red, el servicio de terapia se presta sobre el slice PRIORITY (Gold), configurado para garantizar un acceso preferente a los recursos de red y proteger el tráfico crítico frente a tráfico competitivo. Este enfoque resulta especialmente relevante en entornos reales con coexistencia de múltiples servicios y usuarios.

Las sesiones terapéuticas se desarrollan principalmente en espacios interiores como salas de atención psicológica y salas específicas de terapia, ubicaciones para las que se ha

verificado la disponibilidad de cobertura 5G y capacidad suficiente en las pruebas de validación de red.

5.1.2 Evaluación del rendimiento

Las pruebas de cobertura realizadas en las ubicaciones asociadas al caso de uso de terapia muestran que la red 5G en banda n40 proporciona márgenes de capacidad ampliamente superiores a los requisitos del servicio. En la sala de Atención Psicológica y en la sala de Terapia se alcanzan valores de throughput del orden de 145–159 Mbps en enlace descendente y 28.7 Mbps en enlace ascendente, garantizando conectividad estable incluso en escenarios de uso continuado. En otras zonas interiores cercanas, con condiciones radio menos favorables, se han medido capacidades de hasta 135 Mbps en downlink y 19.5 Mbps en uplink, valores que siguen siendo suficientes para este caso de uso.

En términos de latencia, las pruebas de estabilidad de la red física muestran tiempos de ida y vuelta (RTT) mínimos en torno a 23 ms en ausencia de carga. Bajo condiciones de saturación inducida mediante tráfico sintético, el RTT puede incrementarse significativamente, alcanzando valores del orden de cientos de milisegundos. Estos escenarios representan un caso extremo que no se corresponde con el funcionamiento habitual del servicio terapéutico, pero permiten identificar la latencia como el KPI crítico a proteger mediante mecanismos de calidad de servicio.

Para el escenario de procesamiento delegado de realidad mixta aplicado a terapia, la combinación de latencia de red y tiempo de procesamiento en el MEC resulta determinante. En las configuraciones evaluadas, con resoluciones de trabajo de 1280×480 píxeles y RTT en torno a 50 ms, el sistema alcanza un valor de calidad de experiencia estimada (MOS) aproximado de 4.3. A este valor de RTT se añade un tiempo de inferencia en el nodo MEC del orden de 10–15 ms, manteniendo el retardo extremo a extremo dentro de los márgenes aceptables para la interacción en tiempo real.

En el escenario de regulación emocional, las tasas binarias requeridas son reducidas y la QoE depende principalmente de la latencia de ida y vuelta. Los valores de RTT observados en la red física se sitúan claramente por debajo de los umbrales de perceptibilidad y aceptabilidad definidos para este caso de uso, lo que permite soportar sesiones estables incluso en presencia de múltiples usuarios concurrentes.

5.1.3 Discusión

Los resultados obtenidos confirman que la plataforma desplegada ofrece condiciones adecuadas para la prestación de servicios de terapia en un entorno real. La capacidad disponible en las zonas de uso, junto con la baja latencia observada en condiciones normales de operación, proporciona un margen suficiente frente a los requisitos del servicio.

El uso del slice PRIORITY se identifica como un elemento clave para garantizar la estabilidad del caso de uso, especialmente en situaciones de coexistencia con otros servicios que generan tráfico competitivo. Las pruebas de carga extrema ponen de manifiesto que, sin mecanismos de priorización, la latencia puede degradarse de forma significativa, mientras que la gestión de QoS permite preservar la experiencia de usuario.

Desde el punto de vista de escalabilidad, el escenario de regulación emocional presenta una elevada robustez, al requerir recursos limitados tanto de red como de computación. El procesamiento delegado de realidad mixta, por su parte, constituye el componente más exigente del caso de uso terapéutico y su viabilidad depende de la selección adecuada de algoritmos, resoluciones de trabajo y dimensionado del nodo MEC.

En conjunto, el análisis realizado indica que la plataforma 5G y de edge desplegada es capaz de dar soporte al caso de uso de terapia de forma fiable y estable, estableciendo una base sólida para su explotación continuada y su ampliación futura.

Telepresencia

5.2

5.2.1 Sistemas desplegados

El caso de uso de telepresencia se basa en una aplicación XR de transmisión de vídeo inmersivo en tiempo real, caracterizada por un tráfico dominante en enlace ascendente y una elevada sensibilidad a pérdidas y latencia. El backend de telepresencia se despliega como una función virtualizada en el nodo MEC de la infraestructura 5G privada, encargándose de la gestión de sesiones y de los flujos de datos asociados.

El servicio puede operar en dos modos diferenciados. En el primer modo, el usuario se conecta desde el piso de entrenamiento de la Fundación Juan XXIII, utilizando la red 5G Standalone privada en banda n40. En este entorno controlado, el flujo de telepresencia se transporta sobre el slice PRIORITY (Gold), configurado para proporcionar acceso preferente a los recursos de red y proteger el servicio frente a tráfico competitivo generado por otros usuarios o aplicaciones.

En el segundo modo, el acceso se realiza desde fuera del centro, mediante una conexión remota a través de una red comercial, estableciendo un túnel VPN hacia el backend de telepresencia desplegado en el MEC. En este caso, el tramo de acceso no está bajo el control directo del sistema, por lo que el rendimiento depende de las condiciones de la red pública utilizada.

5.2.2 Evaluación del rendimiento

Las pruebas de cobertura realizadas en el piso de entrenamiento muestran que la red 5G privada proporciona capacidad suficiente para el servicio de telepresencia. En esta ubicación

se han medido valores de throughput del orden de 159 Mbps en enlace descendente y 28.7 Mbps en enlace ascendente, lo que ofrece un margen holgado frente a los bitrates requeridos por la transmisión de vídeo inmersivo, típicamente comprendidos entre 5 y 20 Mbps en uplink.

Las pruebas de slicing con tráfico competitivo confirman que el uso del slice PRIORITY permite preservar la capacidad del flujo de telepresencia. En escenarios con dos equipos de usuario conectados simultáneamente, el slice Gold mantiene valores medios de uplink cercanos a 17 Mbps, mientras que el tráfico no prioritario queda limitado a unos pocos megabits por segundo. Este comportamiento garantiza la estabilidad del servicio incluso en presencia de tráfico concurrente.

En términos de latencia, los tiempos de ida y vuelta observados en la red privada se sitúan en valores bajos en condiciones normales de operación. Aunque las pruebas de carga extrema muestran incrementos significativos del RTT bajo saturación artificial, estos escenarios no representan el funcionamiento habitual del servicio. En condiciones operativas reales, la latencia se mantiene dentro de los márgenes necesarios para una interacción fluida en telepresencia. Como resultado, el valor de calidad de experiencia estimada para este escenario se sitúa en torno a MOS 4.3–4.4, superando el objetivo definido para el caso de uso.

Cuando el acceso se realiza desde fuera del centro mediante una red comercial, no se dispone de medidas directas de cobertura y rendimiento. Para evaluar este escenario se recurre a los resultados obtenidos en red emulada, utilizando configuraciones de microcelda urbana y hotspot representativas de entornos reales.

Los resultados de emulación muestran que, sin mecanismos de priorización, la presencia de carga en la celda provoca pérdidas de paquetes muy elevadas y una degradación severa del servicio, con valores de MOS cercanos a 1. En cambio, cuando se aplica priorización mediante un 5QI adecuado, el sistema es capaz de sostener bitrates de uplink del orden de 6 Mbps, con latencias y pérdidas compatibles con el servicio, alcanzando valores de MOS próximos a 4.3–4.4.

Estos resultados permiten estimar que la telepresencia es viable en entornos de red comercial bajo condiciones favorables de cobertura y congestión moderada. No obstante, la mayor variabilidad de latencia y la ausencia de control sobre el tramo de acceso hacen que la robustez del servicio sea inferior a la observada en el entorno privado.

5.2.3 Discusión

La telepresencia constituye el caso de uso más exigente del proyecto en términos de enlace ascendente y sensibilidad a la calidad de servicio. Los resultados obtenidos demuestran que, en el entorno controlado del piso tutelado, la infraestructura 5G privada desplegada

proporciona capacidad, latencia y estabilidad suficientes para ofrecer el servicio de forma fiable, siempre que se utilicen mecanismos de priorización mediante slicing.

En escenarios de acceso remoto a través de redes comerciales, la viabilidad del servicio depende en gran medida de las condiciones de la red subyacente y de la carga existente. Las estimaciones basadas en red emulada indican que es posible alcanzar niveles aceptables de calidad de experiencia, aunque sin las garantías que ofrece el entorno privado.

En conjunto, el análisis confirma que la plataforma desarrollada cumple los requisitos del caso de uso de telepresencia en entornos controlados y ofrece una base técnica sólida para su extensión a escenarios abiertos, con las limitaciones inherentes a la falta de control sobre la red de acceso.

Teleformación

5.3 5.3.1 Sistemas desplegados

El caso de uso de teleformación constituye un escenario transversal que integra la mayor parte de las tecnologías desarrolladas y validadas en el proyecto. A diferencia de otros casos de uso más acotados, la teleformación combina distintos servicios XR y de procesamiento en función de la fase concreta de la actividad formativa, lo que implica requisitos variables de red y de computación.

La arquitectura se apoya en un backend de teleformación desplegado en el nodo MEC, encargado de la gestión de usuarios, sesiones, contenidos formativos y procesos de evaluación. Sobre este backend se integran distintas funciones virtualizadas, incluyendo una VNF de offloading de algoritmos de inteligencia artificial, que da soporte tanto a tareas de segmentación semántica para escenarios de realidad mixta como a los algoritmos empleados en entrenamiento cognitivo.

Desde el punto de vista de red, la teleformación utiliza la infraestructura 5G Standalone en banda n40, haciendo uso combinado de distintos slices en función del tipo de tráfico. El slice PRIORITY se emplea para flujos críticos en enlace ascendente y servicios sensibles a latencia, mientras que el slice HMDS se utiliza para la distribución de contenidos con tráfico dominante en enlace descendente. El slice DEFAULT queda reservado para tráfico no crítico.

El servicio de teleformación puede desarrollarse en distintos entornos: el piso de entrenamiento o piso tutelado, el centro ocupacional (en espacios como despachos o aulas de informática) y el domicilio del usuario, en este último caso mediante acceso remoto a través de redes comerciales y conexión VPN al backend desplegado en el MEC.

5.3.2 Escenarios funcionales de teleformación

La teleformación se estructura en distintas fases, cada una de las cuales activa un subconjunto específico de tecnologías y presenta requisitos diferenciados.

En la fase de clase teórica, el servicio se basa principalmente en streaming de vídeo inmersivo. El tráfico es dominante en enlace descendente, con tasas binarias moderadas del orden de 15 Mbps por usuario. Los requisitos de latencia son reducidos, siendo prioritarias la estabilidad y la continuidad del servicio, lo que permite escalar a varios usuarios simultáneos utilizando el slice HMDS.

La tutoría individual se apoya en telepresencia XR, con tráfico dominante en enlace ascendente y requisitos más estrictos de latencia y pérdidas. En este escenario, los bitrates típicos en uplink se sitúan entre 5 y 20 Mbps, siendo necesario el uso del slice PRIORITY para garantizar una QoE adecuada, con valores objetivo de MOS superiores a 3.5.

El repaso práctico integra realidad mixta con procesamiento delegado en el MEC. En este caso se combinan flujos de subida y bajada, trabajando con resoluciones de referencia de 1280×480 píxeles a 30 FPS, lo que se traduce en tasas del orden de 10.5 Mbps. La latencia extremo a extremo resulta crítica y depende tanto del RTT de red, típicamente en torno a 50 ms en condiciones favorables, como del tiempo de inferencia en el MEC, del orden de 10–15 ms. Para estas configuraciones se han validado valores de MOS próximos a 4.3, utilizando el slice PRIORITY.

Finalmente, el entrenamiento cognitivo emplea el mismo stack tecnológico que el escenario de regulación emocional, basado en algoritmos de inteligencia artificial de bajo coste computacional ejecutados en el MEC. Este escenario presenta tasas binarias muy reducidas, del orden de 0.1 Mbps tanto en uplink como en downlink, y una elevada tolerancia a la latencia, con umbrales de aceptabilidad de varios cientos de milisegundos. Estas características permiten una elevada escalabilidad, incluso en escenarios con múltiples usuarios concurrentes.

5.3.3 Evaluación del rendimiento

En entornos controlados, como el piso de entrenamiento y el centro ocupacional, las capacidades medidas de la red 5G privada proporcionan un margen suficiente para soportar todas las fases del proceso formativo. En estas ubicaciones se han registrado valores de throughput del orden de 135–159 Mbps en enlace descendente y entre 19.5 y 28.7 Mbps en enlace ascendente, lo que permite acomodar simultáneamente clases teóricas, tutorías individuales y actividades prácticas con procesamiento delegado.

Las pruebas de slicing confirman que el uso del slice PRIORITY permite preservar la capacidad necesaria para los flujos críticos en presencia de tráfico competitivo. En escenarios con múltiples equipos conectados, el slice prioritario mantiene valores de uplink cercanos a 17

Mbps, mientras que el tráfico no crítico queda limitado a unos pocos megabits por segundo. Este comportamiento resulta clave para garantizar la estabilidad de tutorías y repasos prácticos en contextos de uso real.

Cuando la teleformación se realiza desde el domicilio del usuario, el acceso se produce a través de redes comerciales y una conexión VPN hacia el backend del MEC. En este caso, la evaluación se apoya en resultados obtenidos mediante red emulada en escenarios de microcelda urbana y hotspot. Dichos resultados muestran que, sin mecanismos de priorización, la presencia de carga puede provocar pérdidas muy elevadas y una degradación severa del servicio. En cambio, bajo condiciones favorables de cobertura y con gestión adecuada de la calidad de servicio, es posible sostener bitrates de uplink del orden de 6 Mbps, con latencias y pérdidas compatibles con servicios interactivos, alcanzando valores de MOS superiores a 4 en escenarios de tutoría y repaso práctico.

El entrenamiento cognitivo, debido a sus bajos requisitos de red y cómputo, se mantiene como el componente más robusto del caso de uso, siendo viable incluso en condiciones de red menos favorables.

5.3.4 Discusión

La teleformación es el caso de uso más completo e integrador del proyecto, al combinar de forma dinámica streaming XR, telepresencia, procesamiento delegado de realidad mixta y entrenamiento cognitivo sobre una misma plataforma 5G y de edge computing.

Los resultados obtenidos demuestran que, en entornos controlados, la infraestructura desplegada permite soportar de manera fiable todas las fases del proceso formativo, adaptando los recursos de red y cómputo a las necesidades de cada actividad mediante mecanismos de slicing y offloading. En entornos abiertos, como el domicilio del usuario, el servicio sigue siendo viable, aunque condicionado por la variabilidad inherente a las redes comerciales.

En conjunto, el análisis confirma que la plataforma desarrollada cumple los requisitos funcionales y de rendimiento necesarios para la teleformación XR, validando el enfoque del proyecto y su aplicabilidad en escenarios reales y heterogéneos.

6 Apoyo de Fundación Juan XXIII

Este capítulo presenta un análisis de posibles servicios y funcionalidades que podrían desplegarse en distintas zonas de la Fundación Juan XXIII, con el objetivo de ilustrar cómo la infraestructura 5G y la plataforma de edge computing implementadas en el proyecto podrían utilizarse para dar soporte a actividades reales en un centro de estas características. El propósito no es validar nuevos casos de uso, sino identificar opciones de aplicación basadas en las capacidades ya demostradas en los capítulos anteriores.

Para cada servicio propuesto se describe, de forma breve: (i) en qué consiste, (ii) cómo se implementaría apoyándose en los componentes de la plataforma (VR/XR, backend y CMS, offloading de IA, telepresencia y conectividad 5G), (iii) en qué espacios se desarrollaría, y (iv) una valoración preliminar de viabilidad considerando la cobertura medida y el tipo de requisitos esperados.

Servicios en las instalaciones de la Fundación

6.1

6.1.1 Terapias conductuales o cognitivas

Este servicio contempla la realización de sesiones de terapia conductual, estimulación cognitiva y/o entrenamiento en actividades de la vida diaria, apoyadas en entornos de realidad virtual, con el objetivo de trabajar habilidades funcionales y socioemocionales en un entorno controlado, repetible y seguro.

En población con discapacidad intelectual, este tipo de entornos permite abordar algunas cuestiones susceptibles de tratamiento psicológico, como el miedo a las alturas, las escaleras o animales, entre otros, así como entrenamiento en habilidades básicas o instrumentales — como compras, tareas de cocina o rutinas cotidianas— mediante la práctica guiada de actividades que, en el mundo real, pueden implicar riesgo o requerir supervisión continua.

La realidad virtual facilita la modulación progresiva de la dificultad, fomentando un aprendizaje adaptado a los ritmos personales, la repetición de escenarios y la reducción de la carga lingüística, priorizando la interacción visual.

Asimismo, estos entornos pueden emplearse para el entrenamiento de habilidades socio-comunicativas en usuarios con necesidades específicas, permitiendo simular situaciones interpersonales en un contexto estructurado, predecible y controlado.

De este modo, el servicio ofrece un marco flexible para el desarrollo de intervenciones adaptadas a distintos perfiles, integrando la práctica funcional y el entrenamiento social dentro de una misma plataforma tecnológica.

La implementación se basaría en una aplicación VR conectada a un backend de terapia y CMS desplegados en el nodo MEC, complementados con offloading de IA para funcionalidades como segmentación (cuando se integre MR) y algoritmos de baja complejidad para entrenamiento cognitivo. En red, el servicio se cursaría sobre el slice PRIORITY (Gold) para asegurar estabilidad frente a tráfico competitivo, utilizando los mecanismos de control de capacidad validados en la infraestructura.

Las sesiones se desarrollarían principalmente en la sala de terapia en interiores, la localización asociada a terapia (ID 6) presenta valores de referencia de DL ≈ 159 Mbps y UL ≈ 28.7 Mbps, que proporcionan margen suficiente para los flujos de datos típicos de aplicaciones XR.

La viabilidad de este tipo de servicios ya sido estudiada y validada en el proyecto.

6.1.2 Estimulación multisensorial virtual

Este servicio plantea la creación de un entorno virtual de estimulación multisensorial, orientado a la regulación sensorial, la relajación y la exploración guiada. En personas con discapacidad intelectual, este tipo de entornos permite ofrecer estímulos visuales y auditivos controlados, graduables en intensidad y complejidad, adaptándose a distintos perfiles sensoriales y a momentos específicos de la jornada. La virtualización del espacio posibilita además cambiar de escenario de forma inmediata, repetir sesiones con parámetros constantes y reducir la necesidad de acondicionar físicamente salas específicas.

La implementación se basaría en una aplicación de realidad virtual conectada a un backend y CMS desplegados en el nodo MEC, desde el que se gestionan los contenidos y la configuración de las sesiones. El sistema podría apoyarse en offloading de IA para tareas de segmentación o adaptación dinámica del entorno en función de la interacción del usuario, reutilizando el mismo stack tecnológico validado para los escenarios de terapia y entrenamiento cognitivo. Dado que los flujos de datos son moderados y predominantemente en downlink, el servicio podría cursarse sobre slices compartidos, recurriendo a priorización únicamente en escenarios de coexistencia con tráfico más exigente.

Las sesiones se desarrollarían en espacios del centro de día, en zonas destinadas a actividades terapéuticas o de descanso. Aunque no se dispone de medidas específicas de cobertura para todas estas salas, sus características son comparables a las de las salas de terapia evaluadas, donde se han medido capacidades del orden de DL ≈ 145 – 159 Mbps y UL ≈ 28.7 Mbps, suficientes para este tipo de aplicaciones XR.

Desde el punto de vista de viabilidad, el servicio presenta requisitos reducidos de ancho de banda y latencia, con una elevada tolerancia a variaciones en el retardo. A la vista de los resultados de cobertura y estabilidad de la red 5G desplegada, la estimulación multisensorial

virtual resulta compatible con un uso continuado en el centro, sin necesidad de configuraciones específicas más allá de las ya validadas para otros servicios terapéuticos.

6.1.3 Entrenamiento en movilidad

Este servicio está orientado al entrenamiento rutinario en movilidad adaptada con la realización de ejercicios para evitar la sarcopenia, así como mejorar coordinar y equilibrio en personas con discapacidad intelectual, mayores y jóvenes, usando la realidad virtual.

Además, puede utilizarse para entrenar el desplazamiento mediante entornos de realidad virtual, permitiendo practicar de forma segura situaciones que en el entorno real pueden resultar complejas o implicar riesgo, como subir y bajar escaleras, acceder a vehículos, mantener el equilibrio o coordinar movimientos en espacios concurridos.

En personas con discapacidad intelectual, este tipo de entrenamiento facilita la repetición de tareas funcionales y la adquisición progresiva de autonomía, reduciendo la necesidad de supervisión constante y permitiendo adaptar el ritmo y la dificultad a cada usuario. Además, por la experiencia previa en este proyecto, sabemos que les resulta altamente motivante el uso de esta tecnología para desarrollar actividades de su día a día, ofreciendo más oportunidades de participación, no teniendo que depender del apoyo directo del profesional, sino de su supervisión.

La implementación se basaría en una aplicación VR conectada a un backend y CMS desplegados en el nodo MEC, combinada con offloading de IA para tareas de segmentación semántica y análisis del movimiento. De forma complementaria, podrían integrarse sensores para el seguimiento de constantes fisiológicas, cuyos datos se procesan en el edge para su supervisión durante la sesión. Este servicio presenta requisitos moderados de ancho de banda y una mayor sensibilidad a la estabilidad que a la latencia estricta, por lo que se beneficiaría del uso del slice PRIORITY cuando coexiste con otros servicios del centro.

Las sesiones se desarrollarían principalmente en el gimnasio del centro, espacio para el que se han realizado medidas de cobertura. En esta localización se han registrado valores de throughput del orden de DL ≈ 180 Mbps y UL ≈ 28.1 Mbps, proporcionando margen suficiente para soportar aplicaciones XR y la transmisión de datos adicionales procedentes de sensores.

A la vista de los resultados de validación de la red y del sistema de edge, este servicio se considera viable en el entorno del gimnasio. La combinación de conectividad 5G estable y procesamiento en el MEC permite soportar sesiones continuadas, con capacidad suficiente para integrar análisis de movimiento y monitorización básica sin comprometer la experiencia de usuario.

6.1.4 Formación

Este servicio contempla la utilización de la plataforma desarrollada para dar soporte a actividades formativas vinculadas a tareas reales que se llevan a cabo en distintas áreas de la Fundación Juan XXIII. A diferencia del caso de teleformación analizado en el capítulo anterior, el foco aquí no está en la validación técnica del servicio, sino en su adaptación a contextos formativos concretos, en los que la formación se integra directamente en el entorno donde se realiza la actividad.

La implementación reutilizaría el mismo stack tecnológico descrito para la teleformación, combinando aplicaciones XR, backend y CMS desplegados en el MEC, y, cuando sea necesario, funciones de telepresencia y offloading de IA para segmentación o apoyo cognitivo. En función del tipo de actividad, se activarían distintos modos de uso: contenidos formativos en vídeo para fases más teóricas, asistencia remota mediante telepresencia para tutorías o supervisión, y soporte en realidad mixta para el repaso práctico de tareas específicas.

Las actividades formativas se desarrollarían en los espacios donde tienen lugar las tareas reales: la cocina del piso de entrenamiento para formación en hostelería, la terraza de la cafetería para actividades relacionadas con soluciones verdes o jardinería, y el taller ocupacional (zona 11) para procesos como la elaboración de bolsas de papel (“Paper lovers”). Este enfoque permite que la formación se realice en un contexto funcional y familiar para los usuarios.

Desde el punto de vista de viabilidad, y a la luz de los resultados presentados en el capítulo 5, la infraestructura 5G desplegada proporciona cobertura y capacidad suficientes en estas zonas para soportar los distintos modos de formación. El uso de slicing y edge computing permite adaptar dinámicamente los recursos de red y cómputo a cada actividad, facilitando la prestación del servicio en entornos heterogéneos dentro de la Fundación.

6.1.5 Asistencia a eventos y formaciones

Este servicio se orienta a facilitar la asistencia a eventos y formaciones mediante soluciones de telepresencia, permitiendo que personas que se encuentren fuera del centro, o que presenten dificultades para el contacto personal cercano, puedan acceder a actividades que se desarrollan de forma presencial. Este enfoque resulta especialmente relevante para usuarios con necesidades específicas de interacción social, así como para situaciones en las que la presencia física no es posible.

La implementación se basaría en un sistema de telepresencia apoyado en el backend desplegado en el nodo MEC, encargado de la gestión de sesiones y flujos de vídeo. El acceso de los usuarios remotos se realizaría mediante una conexión VPN al backend, reutilizando el mismo esquema validado para los escenarios de telepresencia. Dado que el tráfico es predominantemente en enlace descendente y con requisitos moderados de interacción, el

servicio puede cursarse sobre slices compartidos, recurriendo a priorización únicamente cuando sea necesario garantizar estabilidad frente a tráfico concurrente.

Los eventos y formaciones se desarrollarían principalmente en el auditorio del centro, desde donde se emitiría el contenido hacia usuarios remotos. En esta localización se han medido valores de throughput del orden de DL ≈ 103 Mbps y UL ≈ 19.5 Mbps, suficientes para soportar la transmisión de vídeo y audio asociados a este tipo de servicio.

A la vista de los resultados de cobertura y estabilidad de la red 5G desplegada, la asistencia remota a eventos se considera viable, con un comportamiento robusto incluso en presencia de varios usuarios conectados simultáneamente, y sin requerir configuraciones específicas más allá de las ya validadas para otros servicios de telepresencia.

6.1.6 Vigilancia de exteriores

Este servicio contempla el uso de soluciones de telepresencia y transmisión de vídeo para la vigilancia y supervisión de zonas exteriores de la Fundación, como accesos, áreas de aparcamiento o zonas logísticas. El objetivo es facilitar el control remoto de estos espacios, mejorando la seguridad y permitiendo una supervisión continua sin necesidad de presencia física constante.

La implementación se basaría en sistemas de captura de vídeo conectados al backend de telepresencia desplegado en el nodo MEC, desde el que se gestionarían los flujos y el acceso a las imágenes en tiempo real. Este servicio presenta requisitos moderados de ancho de banda y baja sensibilidad a la latencia, por lo que puede cursarse sobre slices compartidos, recurriendo a mecanismos de priorización únicamente en escenarios de coexistencia con servicios más críticos.

La vigilancia se desarrollaría en zonas exteriores del centro, como el acceso al parking, la playa de camiones o áreas de tránsito. En estas ubicaciones se han realizado medidas de cobertura que muestran valores de throughput del orden de DL ≈ 106 – 116 Mbps y UL ≈ 20 – 22 Mbps, lo que proporciona margen suficiente para la transmisión de vídeo en tiempo real desde varios puntos de captura.

De acuerdo con los resultados de las pruebas de cobertura y estabilidad, la red 5G desplegada ofrece conectividad fiable en exteriores, lo que permite considerar viable la implantación de este servicio como complemento a los sistemas de vigilancia existentes, sin requerir configuraciones específicas adicionales.

Servicios en remoto

Además de los servicios que pueden prestarse en las instalaciones de la Fundación, la plataforma desarrollada permite extender determinadas funcionalidades a usuarios

remotos, mediante el acceso a los backends desplegados en el centro a través de una conexión VPN sobre redes comerciales. Este enfoque posibilita ofrecer continuidad en la prestación de servicios a personas que se encuentren en su domicilio u otros entornos externos, sin necesidad de replicar la infraestructura de edge computing fuera de la Fundación.

Los servicios que pueden prestarse en remoto se basan fundamentalmente en los escenarios de telepresencia y teleformación analizados en el capítulo anterior. En este contexto, es posible dar soporte a tutorías individuales, seguimiento formativo, asistencia a eventos o acceso a contenidos formativos, utilizando aplicaciones XR o soluciones de vídeo inmersivo conectadas al backend del MEC. Asimismo, servicios con requisitos reducidos de red, como el entrenamiento cognitivo, resultan especialmente adecuados para su prestación en entornos no controlados.

Tal y como se ha discutido previamente, la viabilidad de estos servicios en remoto depende de las condiciones de la red de acceso comercial, en particular de la cobertura y la congestión existentes. No obstante, los resultados obtenidos en red emulada indican que, bajo condiciones favorables, es posible alcanzar niveles de calidad de experiencia compatibles con los objetivos del proyecto, permitiendo ampliar el alcance de la plataforma más allá de las instalaciones físicas de la Fundación.

7 Conclusiones

Los resultados obtenidos muestran que la red 5G desplegada presenta un comportamiento estable en condiciones normales de operación, identificándose la latencia como el KPI más crítico en escenarios de saturación. Los modelos de QoE validados permiten estimar de forma consistente la experiencia de usuario en los distintos escenarios XR considerados, facilitando la evaluación conjunta de red y servicio.

El uso de edge computing y técnicas de offloading de IA se confirma como un elemento clave para soportar servicios interactivos con requisitos estrictos. En este contexto, el caso de uso de terapia presenta requisitos moderados y una elevada robustez en entornos controlados, mientras que la telepresencia se identifica como el escenario más exigente en enlace ascendente y el más dependiente de mecanismos de priorización mediante slicing. La teleformación actúa como un caso integrador, combinando múltiples tecnologías y patrones de tráfico sobre una misma plataforma.

De forma global, los servicios analizados resultan plenamente viables en entornos controlados y parcialmente viables en entornos abiertos, condicionados por las características de la red de acceso. La conexión remota mediante VPN permite extender determinados servicios más allá de las instalaciones del centro, ampliando el alcance funcional de la plataforma. Asimismo, la solución demuestra flexibilidad para adaptarse a distintos espacios y actividades de la Fundación Juan XXIII.

En conjunto, los resultados confirman la idoneidad de la plataforma desarrollada para aplicaciones XR en contextos sociosanitarios reales, y sientan una base sólida para futuras ampliaciones y para la explotación del sistema en escenarios de uso continuado.

8 Acrónimos

5G	Fifth Generation (quinta generación de redes móviles)
5QI	5G Quality Indicator
AI / IA	Artificial Intelligence / Inteligencia Artificial
API	Application Programming Interface
CMS	Content Management System
DL	Downlink
FPS	Frames Per Second
gNB	gNodeB
GPU	Graphics Processing Unit
HMD	Head Mounted Display
KPI	Key Performance Indicator
MEC	Multi-access Edge Computing
MOS	Mean Opinion Score
MR	Mixed Reality
QoE	Quality of Experience
QoS	Quality of Service
RTT	Round Trip Time
SA	Standalone
SVM	Support Vector Machine
TDD	Time Division Duplex
UE	User Equipment
UL	Uplink
VNF	Virtual Network Function



VPN	Virtual Private Network
VR	Virtual Reality
XR	Extended Reality

9 Índice de tablas

Tabla 1. Resumen de resultados de las pruebas de cobertura	10
Tabla 2. Medidas de tráfico de uplink para distintas configuraciones de slice	12
Tabla 3. Medidas de tráfico de downlink para distintas configuraciones de slice	13
Tabla 4. Medidas de tráfico de uplink para la configuración con dos módems	14
Tabla 5. Escenarios de simulación y requisitos.....	32
Tabla 6. Parámetros de simulación del escenario de teleportación virtual en red emulada.	33
Tabla 7. MOS estimado para el escenario de teleportación virtual bajo distintas configuraciones de carga y 5QI.....	35
Tabla 8. Parámetros de simulación del escenario de procesamiento delegado	35
Tabla 9. MOS estimado para el escenario de procesamiento delegado en función de la banda y el número de usuarios	38
Tabla 10. Parámetros de simulación del escenario de regulación emocional.....	39
Tabla 11. MOS estimado para el escenario de regulación emocional	41
Tabla 12. Configuración de slices de red	42
Tabla 13. Configuración de la prueba	42
Tabla 14. Configuración de la prueba del Escenario 1	42
Tabla 15. Resumen del Escenario 1, sin slicing, con tráfico competitivo de 20 Mbps. 44	
Tabla 16. Resumen del Escenario 1, con slicing, con tráfico competitivo de 6Mbps... 46	
Tabla 17. Configuración de la prueba del Escenario 2.....	50
Tabla 18. Resumen del Escenario 2, con slicing, con tráfico competitivo de 100 Mbps de downlink	52
Tabla 19. Resumen del Escenario 3, con slicing, con tráfico competitivo de 100 Mbps de bandwidth.....	55
Tabla 20. Requisitos y valores de referencia para las pruebas de carga de los escenarios de computación en el edge	69

Tabla 21. Resumen del Escenario 2: métricas de inferencia, tasa de cuadros y uso de memoria GPU	78
---	----

10 Índice de figuras

Fig. 2.1. Localización de las pruebas de cobertura en la Fundación Juan XXIII	9
Fig. 2.2. Monitorización del tráfico de red durante las pruebas de cobertura	11
Fig. 2.3. Monitorización del tráfico de uplink con cambio de configuración de slices .	12
Fig. 2.4. Monitorización del tráfico de downlink con cambio de configuración de slices	13
Fig. 2.5. Monitorización del tráfico de uplink para la configuración con dos módems	14
Fig. 2.6. Detalle de la configuración de carga en las pruebas de estabilidad.....	15
Fig. 2.7. Monitorización del tráfico en el core (UL/DL, dos slices) durante las pruebas de estabilidad	16
Fig. 2.8. Monitorización del tiempo de ping (RTT) desde un módem durante las pruebas de estabilidad	16
Fig. 2.9. Monitorización de la carga de CPU del core durante las pruebas de estabilidad	17
Fig. 2.10. Monitorización del consumo de energía en la banda base (verde) y RRH (amarillo) durante las pruebas de estabilidad	18
Fig. 2.11. Monitorización de la Calidad de señal del módem en posición 5 (aula de informática) durante las pruebas de estabilidad.....	18
Fig. 2.12. Monitorización de la calidad de señal del módem en posición 1 (despachos tercera planta) durante las pruebas de estabilidad.....	18
Fig. 2.13. Detalle del tráfico en el core (UL/DL, dos slices) durante las pruebas de estabilidad	19
Fig. 2.14. Detalle del tiempo de ping (RTT) desde un módem durante las pruebas de estabilidad	19
Fig. 2.15. Detalle de la calidad de señal del módem en posición 5 (aula de informática) durante las pruebas de estabilidad	20
Fig. 2.16. Detalle de la calidad de señal del módem en posición 1 (despachos tercera planta) durante las pruebas de estabilidad	20
Fig. 2.17. Detalle del consumo de energía en banda base (verde) y RRH (amarillo) durante las pruebas de estabilidad	21

Fig. 3.1. Descripción general de la lógica del modelo modular de QoE.....	23
Fig. 3.2. Relación entre μ y MOS para diferentes valores de σ	24
Fig. 3.3. Curvas de MOS para el escenario de teleportación virtual	25
Fig. 3.4. Curvas de MOS para el escenario de procesamiento delegado	26
Fig. 3.5. Tiempo de inferencia de referencia para cada resolución	26
Fig. 3.6. Curva de MOS para el escenario de regulación emocional	27
Fig. 3.7. Parámetro de desvanecimiento en microcelda urbana (UMi): modelo frente a medidas	29
Fig. 3.8. Distribución de los valores de sombreado obtenidos en una campaña de macroceldas rural (RMa).....	29
Fig. 3.9. Distribución de SINR: medida y estimada por ambas versiones de FikoRE....	30
Fig. 3.10. Distribución de tráfico alcanzado en entorno rural: real vs. emulado	31
Fig. 3.11. Trayectorias e histogramas de SINR para una simulación con 10 usuarios.	34
Fig. 3.12. Throughput y latencia de uplink del usuario emulado en función del número de usuarios concurrentes.....	34
Fig. 3.13. Trayectorias en banda n40 (izquierda) y n78 (derecha).....	36
Fig. 3.14. Histogramas de SINR en banda n40 (izquierda) y n78 (derecha).....	37
Fig. 3.15. Gráficas de KPIs de aplicación para los distintos escenarios de computación delegada	38
Fig. 3.16. Distribución de UEs e histogramas de SINR para el escenario de regulación emocional	40
Fig. 3.17. KPIs para el escenario de regulación emocional.....	41
Fig. 3.18. Tráfico de uplink recibido en el core desde los UEs A (rojo) y B (verde)	43
Fig. 3.19. Medidas en el backend: bit rate de video (recibido), tasa de pérdida de paquetes y retardo medio en de paquete (one-way).....	44
Fig. 3.20. Tráfico de uplink recibido en el core desde los UEs A (rojo) y B (verde)	46
Fig. 3.21. Medidas en el backend: bit rate de video (recibido), tasa de pérdida de paquetes y retardo medio en de paquete (one-way).....	46

Fig. 3.22. Detalle del uplink de los UES A (rojo, escenario de prueba) y B (verde, ruido) para el caso de prueba de 20 Mbps de bit rate de video. Sin slicing en la parte superior, con slicing en la inferior.....	48
Fig. 3.23. Comparación del bit rate recibido sin (izquierda) o con (derecha) el slicing activo	48
Fig. 3.24. Comparación del la tasa de pérdidas sin (izquierda) o con (derecha) el slicing activo	49
Fig. 3.25. Comparación del la latencia sin (izquierda) o con (derecha) el slicing activo	49
Fig. 3.26. MOS estimado en función del bitrate transmitido desde A, asumiendo que desde B se genera simultáneamente 20 Mbps	50
Fig. 3.27. Tráfico de downlink enviado desde el core a los UEs A (morado) y B (amarillo)	51
Fig. 3.28. Medidas en el backend: bit rate de video (recibido), tasa de paquetes por segundo (Frame Rate) y retardo medio en de paquete (one-way)	52
Fig. 3.29. MOS estimado en función de la resolución	53
Fig. 3.30. Tráfico de downlink enviado desde el core a los UEs A (morado) y B (amarillo)	54
Fig. 3.31. Medidas en el backend: bit rate de datos (recibido), tasa de paquetes por segundo (Sample Rate) y retardo medio en de paquete (one-way)	55
Fig. 4.1. Arquitectura general del Sistema Edge	60
Fig. 4.2. Consola de Proxmox para el servidor OpenEdge MEC <i>Cabazon</i>	61
Fig. 4.3. Consola de Proxmox para el servidor OpenEdge MEC <i>Casals</i>	62
Fig. 4.4. Monitorización de una de las VNFs de <i>Casals</i> (core 5G) desde la consola Proxmox	63
Fig. 4.5. Monitorización de una de las VNFs de <i>Cabazon</i> (<i>Clemente10</i>) desde la consola Proxmox	64
Fig. 4.6. Imágenes docker instaladas en <i>Clemente10</i> para ejecutar los distintos algoritmos de offloading de IA	65
Fig. 4.7. Monitorización de la GPU con Nvidia-SMI	65
Fig. 4.8. Monitorización de la CPU con htop	66

Fig. 4.9. Imágenes docker ejecutando en <i>Clemente11</i> para dar servicio a los casos de uso	66
Fig. 4.10. Imágenes docker ejecutando en <i>Pio12</i> (backend de telepresencia)	67
Fig. 4.11. Imágenes docker ejecutando en <i>Leon13</i> (backend de teleformación / terapia)	67
Fig. 4.12. Imágenes docker ejecutando en <i>Gregorio15</i> (CMS para teleformación / terapia)	67
Fig. 4.13. Servicios del Stack TIG configurados en <i>Benedicto16</i> mediante systemd ..	68
Fig. 4.14. Tiempo de inferencia para el algoritmo Thundernet con resolución 1280x480 ..	71
Fig. 4.15. Tiempo de inferencia para el algoritmo Thundernet con resolución 2540x960 ..	71
4.2.1.3 Fig. 4.16. Tiempo de inferencia para el algoritmo Yolo + EgoHos con resolución 1280x480	72
Fig. 4.17. Tiempo de inferencia para el algoritmo Yolo + EgoHos con resolución 2540 x 960	72
Fig. 4.18. Tiempo de inferencia para el algoritmo Yolo +Thundernet con resolución 1280x480	73
Fig. 4.19. Tiempo de inferencia para el algoritmo Yolo + Thundernet con resolución 2540x960	74
Fig. 4.20. Tiempo de inferencia para el algoritmo Yolo con resolución 1280x480.....	74
Fig. 4.21. Tiempo de inferencia para el algoritmo Thundernet con resolución 2540x960 ..	75
Fig. 4.22. Tiempo de inferencia para el algoritmo Yolo + Chroma con resolución 1280x480	75
Fig. 4.23. Tiempo de inferencia para el algoritmo Yolo + Chroma con resolución 2540x960	76
Fig. 4.24. Tiempo de inferencia para el algoritmo Chroma con resolución 1280x480 ..	76
Fig. 4.25. Tiempo de inferencia para el algoritmo Chroma con resolución 2540x960 ..	76
4.2.1.9 Fig. 4.26. Tiempo de inferencia para el algoritmo Adaptive Chroma con resolución 1280x480	77
Fig. 4.27. Tiempo de inferencia para el algoritmo Adaptive Chroma con resolución 2540x960	77

Fig. 4.28. Tiempo de inferencia para el algoritmo SVM lineal para regulación emocional 79

11 Referencias

Díaz, C., Pérez, P., Cabrera, J., Ruiz, J. J., & García, N. (2020). XLR (pixel loss rate): A lightweight indicator to measure video QoE in IP networks. *IEEE Transactions on Network and Service Management*, 17(2), 1096-1109.

Cortés, C., Pérez, P., & García, N. (2023). Understanding latency and qoe in social xr. *IEEE Consumer Electronics Magazine*, 13(3), 61-72.

Braithwaite, J. J., Watson, D. G., Jones, R., & Rowe, M. (2013). A guide for analysing electrodermal activity (EDA) & skin conductance responses (SCRs) for psychological experiments. *Psychophysiology*, 49(1), 1017-1034.

Benedek, M., & Kaernbach, C. (2010). A continuous measure of phasic electrodermal activity. *Journal of neuroscience methods*, 190(1), 80-91.

González-Morín, D., Lopez-Morales, M. J., Perez, P., Armada, A. G., & Villegas, A. (2023). FikoRE: 5G and beyond RAN emulator for application level experimentation and prototyping. *IEEE Network*, 37(4), 48-55.